



JO HERMANS

With illustrations by Wiebke Drenckhan

Physics in Daily Life

Foreword by Sir Arnold Wolfendale



Physics in Daily Life

ABOUT THE AUTHORS

Prof. L.J.F. Hermans is Emeritus Professor of Physics at Leiden University, The Netherlands. In addition to his academic teaching and research career he was quite active in promoting and explaining science for the general public. In this context he published, among others, a book about Every-day science (in Dutch) and two books about Energy (in Dutch and English). He is presently Science Editor of Europhysics News. He was appointed Knight in the Order of Oranje Nassau by Queen Beatrix in 2010.

Wiebke Drenckhan is CNRS researcher at the *Laboratoire de Physique des Solides* at the outskirts of Paris, where she tries to unravel the physical properties of soft materials, such as foams or emulsions. In her spare time she finds great pleasure in letting scientific issues come to life with pen and paper in the form of illustrations or cartoons.



Physics in Daily Life

JO HERMANS

With illustrations by Wiebke Drenckhan



17, avenue du Hoggar – P.A. de Courtabœuf
BP 112, 91944 Les Ulis Cedex A

This is a collection of 'Physics in Daily Life' columns which appeared
in Europhysics News, volumes 34 - 42 (2003 – 2011)

Mise en pages : Patrick Leleux PAO

Imprimé en France
ISBN : 978-2-7598-0705-5

Tous droits de traduction, d'adaptation et de reproduction par tous procédés, réservés pour tous pays. La loi du 11 mars 1957 n'autorisant, aux termes des alinéas 2 et 3 de l'article 41, d'une part, que les «-copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective-», et d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation intégrale, ou partielle, faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (alinéa 1^{er} de l'article 40). Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles 425 et suivants du code pénal.

© EDP Sciences, 2012

CONTENTS

<i>Foreword</i>	7
1. The human engine.....	11
2. Moving around efficiently	14
3. Hear, hear.....	16
4. Drag'n roll.....	19
5. Old ears.....	22
6. Fresh air	25
7. Diffraction-limited photography	28
8. Time and money	31
9. Blue skies, blue seas	33
10. Cycling in the wind	36
11. Seeing under water.....	39
12. Cycling really fast.....	41
13. Water from heaven	43
14. Surviving the sauna	45
15. Black vs. white	48
16. Hearing the curtains	50
17. Fun with the setting sun.....	52
18. NOT seeing the light.....	54

19. Thirsty passengers	57
20. The sauna – revisited	59
21. Refueling	62
22. Counting flames	64
23. Drink or drive	66
24. Feeling hot, feeling cold	68
25. The way we walk	70
26. Wine temperature	72
27. Over the rainbow	74
28. New light	77
29. Windmill nuisance	80
30. Fog and raindrops	83
31. Why planes fly	85
32. Heating problems	87
33. Bubbles and balloons	89
34. Funny microwaves	92
35. Brave ducks	95
36. Muddy cyclist	98
37. Flying (s)low	100
38. Funny ice	103
39. Amazing candle flames	106
40. Capricious suntime	109

FOREWORD

The history of Physics in Europe is one of brilliance and the sun is still shining, indeed it is getting ever brighter, despite the economic problems. The European Physical Society is a composite of all the national physical societies and it occupies an important role in providing advice to its members and a forum for discussion.

Its house journal, *Europhysics News*, is an exciting small publication, packed with interesting articles about conferences, national societies, highlights from European journals and ‘features’. In addition there has been, for the past decade, a page entitled ‘Physics in Daily Life’. The present volume is a collection of these pages and is a feast of erudition and humour, by way of the excellent accompanying cartoons as well as the subject matter.

It is easy for those of us steeped in our disciplines, of astrophysics, condensed matter, nuclear physics, or whatever, to think that ‘everyday physics’ is child’s play compared with the deep subtleties of our chosen subjects. Surely, if we can understand the mysteries of parallel universes, the behaviour of superconductors or exotic atomic nuclei, the V-shaped pattern of a duck’s wake in the lake at the local Wildfowl Park will be a ‘piece of cake’. However, it would be wise, before telling ones child/grandchild/lady or gentleman friend or... to read the contribution ‘Brave Ducks’ herein. Quite fascinating...

In a similar vein, the Astrophysicist who knows all about the recently found bubbles in the interstellar medium just outside the heliopause, and the Local Bubble in which the solar system is immersed, had better read the 'Bubbles and Balloons' piece before setting himself or herself up as an authority on such matters at the next Christmas Children's Party.

Michael Faraday, that physicist of genius, whose discoveries led to the electrical power industry amongst many other things, lectured for one hour on the physics and chemistry of the candle flame. He probably knew the points made in 'Amazing Candle Flames' (contribution number 39) but I didn't. Henceforth, my over-dinner description of the candle flames at the table will be the envy of my guests – even the physicists and chemists amongst them (unless they happen to belong to the EPS).

Turning to our activities on the high seas, where many of us use our SKI funds ('Spending the kids' inheritance') to take exotic cruises, we have the oft-sought 'green flash' from the sun as it sinks below the horizon. Wearing our tuxedos and leaning over the rail with our new-found friends, we have languidly explained what we should have seen as the sun gently disappeared (only occasionally does it make an appearance). Beware, however, your explanation may not be quite right – 'Fun with the setting sun' (contribution number 17) will put you right. Even one's description of why the sea sometimes looks blue may turn out to have been wrong! Better to take with you an absorption curve for water, from 400-700 nm, to nonchalantly fish out of your pocket at the appropriate moment.

Now to taxi-drivers, most are sources of information, freely imparted, and their views are strongly held. In order to keep one step ahead it would be wise to dip into our compendium and produce such gems as 'Hearing the Curtain' (contribution number 16) which relates to the reason why we all like to sing in the bath. The driver will be enthralled when you explain that the sound absorption properties of the curtains are the same whether they are drawn shut

or quite open. Indeed it may lead to some interesting descriptions of sights that the taxi driver himself has witnessed during his late night excursions.

So, what about this collection? For me, at least, it scores 10/10 and I recommend it to all who have an interest in the physical world and explanations of what seem to be – but are often not – simple phenomena. Not only that, but buy it for your friends and relatives.

Arnold Wolfendale

(Sir Arnold Wolfendale FRS is a Past-President of the EPS. He is emeritus Professor of Physics in Durham University, UK)



© David Haldane.

This page intentionally left blank

1

The human engine (and how to keep it cool)



We don't usually think of ourselves in that way, but each of us is an engine, running on sustainable energy. It differs from ordinary engines in more than just the fuel. The human engine cannot be shut off; for instance, it keeps idling even if no work is required. This is needed to keep the system going, to keep our heart pumping, for example, and to keep the temperature around 37 °C. Because – and here is another difference – our human engine works in a very small temperature range.

It's interesting to look at this a bit more quantitatively. Our daily food has an energy content of 8 to 10 MJ. That, incidentally, is equivalent to a quarter of a litre of gasoline, barely enough to keep our car going on the highway for about 2 minutes. Those 8 to 10 MJ per day represent just about 100 W on a continuous basis. Only a small fraction is needed to keep our heart pumping, as we can easily estimate from a $p\Delta V$ consideration (p being on the order of 10 kPa and ΔV on the order of 0.1 litre, with a heart beat frequency of around 1 Hz).

In the end, those 100 W are released as heat: by radiation, conduction and evaporation. Under normal conditions, sitting behind our desk in our usual clothing in an office at 20 °C, radiation and conduction are the leading terms, while evaporation gives only a small contribution. But when we start doing external work, on a home trainer, for example, the energy consumption goes up, and so does the heat production. Schematically, the total energy consumption P_{tot} vs. external work P_{work} is shown in the figure, where an efficiency of 25% has been assumed. Thus, if we work with a power of 100 W, we increase the total power by 400 W, and the heat part P_{heat} by 300 W.

Now our body must try to keep its temperature constant. That's not trivial: if we don't change clothing, or switch on a fan to make the temperature gradients near our skin somewhat larger, the radiation and conduction terms cannot change much. They are determined by the difference between the temperature of our skin and clothing on the one hand, and the ambient temperature on the other. When working hard, we increase that difference only slightly. Granted, due to the enhanced blood circulation, our skin temperature will get closer to that of our inner body, but the limit is reached at 37 °C.

Fortunately, there is also the evaporation term. Sweating comes to our rescue, as also, of course, does drinking! Each additional 100 W of released heat that has to be compensated by evaporation requires a glass of water per hour (0.15 litre, to be more precise). The various terms are schematically shown in the figure.

One conclusion: heavy exercise requires evaporation. Don't try to swim a 1000 m world record if your pool is heated to 37 °C. You might not live to collect your prize, because where would the heat go?

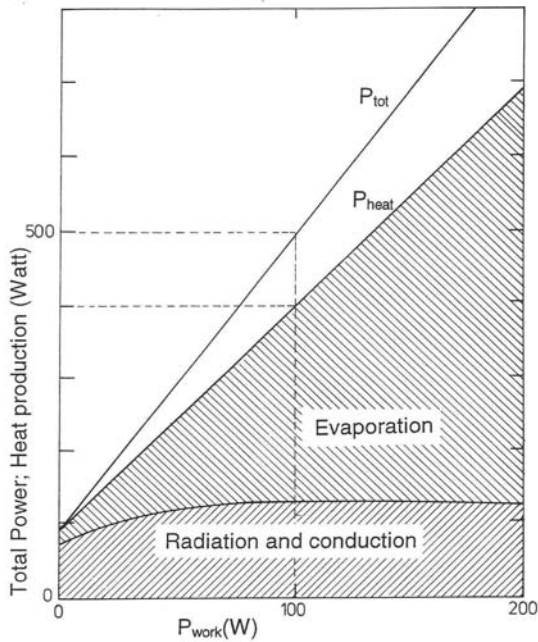


Image 1.1 | Total energy production, heat production and heat release vs. external mechanical power, schematically.

2

Moving around efficiently



Ever considered the efficiency of a human being moving from A to B? Not by using a car or a plane, but just our muscles. Not burning oil, but food.

Many physicists will immediately shout: A bike! Use a bicycle! It is because we all know from experience that using wheels gets us around about five times as fast as going by foot with the same effort.

But just how efficient is a bike ride? First, we have to examine the human engine. The power we produce is easily estimated by climbing stairs. If we want to do that on a more or less continuous basis, one step per second is a reasonable guess. Assuming a step height of 15 cm

and a mass of 70 kg, this yields a power of roughly 100 W. Mountain climbers will find the assumed vertical speed quite realistic, since it takes us about 500 m high in an hour, and that is pretty tough exercise.

Riding our bike is pretty much like climbing the stairs: same muscles, same pace. In other words, we propel our bike with about 100 W of power. But that is not the whole story. The efficiency of our muscles comes into play. For this type of activity, the efficiency is not so bad (a lot better than *e.g.* weight lifting). We may reach 25%. The total energy consumption needed for riding is therefore around 400 W.

What does this tell us about the overall transport efficiency? How does this compare with other vehicles? Now it's time to do a back-of-the-envelope calculation. If we express 400 W of continuous energy use in terms of oil consumption per day, we find pretty much exactly one litre per day, given that the heat of combustion for most types of oil and gasoline is about 35 MJ per litre. In other words: if, for the sake of the argument, we ride for 24 hours continuously without getting off our bike, we have used the equivalent of 1 litre of gasoline for keeping moving. How far will that get us? That, of course, depends on the type of bike, the shape of the rider, and other parameters. If we take a speed of 20 km/h as a fair estimate, the 24 hours of pedaling will get us as far as 480 km. In other words: a cyclist averages about 500 km per litre.

That's not bad, compared to a car, or even a motorbike. So, we should all ride our bike if we want to conserve energy? Careful, there is a catch. We have been moving on food, not gasoline or oil. And it takes a lot more energy to get our food on the table than its energy content may suggest. A glass of milk, for example, takes roughly 0.1 litre of oil, and a kg of cheese even about 1 litre. It's because the cow has to be milked, the milk has to be cooled, transported, heated, bottled, cooled again, transported again, etc. It's the same (or worse) for cheese, meat, etc.

Conclusion: Riding our bike is fun. It's healthy. It keeps us in good shape. And, if we have to slim down anyway, it conserves energy. Otherwise – I hate to admit it – a light motorbike, if not ridden too fast, might beat them all.

3

Hear, hear



Even a tiny cricket can make a lot of noise, without having to ‘refuel’ every other minute. It illustrates what we physicists have known all along: audible sound waves carry very little energy. Or, if you wish, the human ear is pretty sensitive – if the sound waves are in the right frequency range, of course.

Exactly how our ears respond to sound waves has been sorted out by our biophysical and medical colleagues, and is illustrated by the familiar isophone plots that many of us remember from the textbooks. They are reproduced here for convenience.

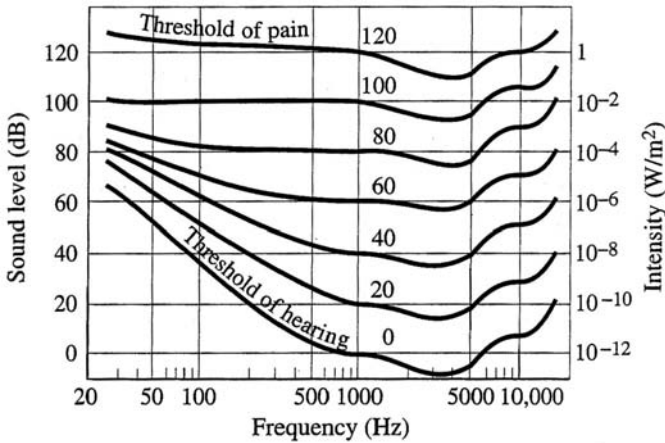


Image 3.1 | Isophone curves, with vertical scales in dB (left) and W/m² (right).

Each isophone curve represents sound that seems to be equally loud for the average person.

The figure reminds us that the human ear is not only rather sensitive, but that it also has an astonishingly large range: 12 orders of magnitude around 1 kHz. This is, in a way, a crazy result, if we think of noise pollution. It means that, if we experience noise loud enough to reach the threshold of pain, and we assume that the sound intensity decays with distance as $1/r^2$, we would have to increase the distance from the source r by a factor of 10^6 to get rid of the noise. Or, if we stand at 10 m from the source, we would have to walk away some 10 000 km.

Here we have assumed that the attenuation can be neglected, since we have been taught that sound wave propagation is an adiabatic process. Obviously, real life isn't that simple. There are several dissipative terms. For example, think of the irreversible heat leaks between the compressed and the expanded air. An interesting feature here is that the classical absorption coefficient is proportional to the frequency squared, which makes distant thunder rumble. Then there is attenuation by obstacles. In addition, there is the curvature of the earth, and the curvature of the sound waves themselves, usually away from the earth due to the vertical temperature gradient. Without loss terms like these, forget a solid sleep.

A second feature worth noticing is the *shape* of the curves. Whereas the pain threshold curve is relatively flat, the threshold of hearing increases steeply with decreasing frequency below 1 kHz. If we turn our audio amplifier from a high to a low volume, we tend to lose the lowest frequencies. The 'loudness control' is intended to compensate for this.

Finally, it is interesting to notice the *magnitude* of the sound intensity. How much sound energy do we produce when we speak? Let us assume that the listener hears us speak at an average sound level of 60 dB, which corresponds to 10^{-6} W/m² as seen from the right-hand vertical scale. Assuming that the listener is at 2 m, the energy is 'smeared out' over some 10 m². This means that we produce, typically, 10^{-5} W of sound energy when we talk. That is very little indeed. During our whole life, even if we talk day and night and we get to live 100 years, we will not talk for more than 10^6 hours. With the above 10^{-5} W, this means a total energy of 10 Wh. Even at a relatively high price of € 0.50/kWh, this boils down to less than one cent for life-long speaking. Cheap talk, so to speak.

4

Drag'n roll



Whether we ride our bike or drive our car, there is resistance to be overcome, even on a flat road; that much we know. But when it comes to the details, it's not that trivial. Both components of the resistance – rolling resistance and drag – deserve a closer look. Let's first remember the main cause of the rolling resistance. It's not friction in the ball bearings, provided they are well greased and in good shape. It's the tires, getting deformed by the road. In a way, that may be surprising: the deformation seems elastic, it's not permanent. But there is a catch here: the forces for compression are not compensated for by those for expansion of the rubber

(there is some hysteresis, if you wish). The net work done shows up as heat.

The corresponding rolling resistance is, to a reasonable approximation, independent of speed (which will become obvious below). It is proportional to the weight of the car, and is therefore written: $F_{\text{roll}} = C_r mg$, with C_r the appropriate coefficient. Now we can make an educated guess as to the value of C_r . Could it be 0.1? No way: this would mean that it would take a slope of 10% to get our car moving. We know from experience that a 1% slope would be a better guess. Right! For most tires inflated to the recommended pressure, $C_r = 0.01$ is a standard value. By the way: for bicycle tires, with pressures about twice as high, C_r can get as low as 0.005.

The conclusion is that, for a 1000 kg car, the rolling resistance is about 100N.

What about the drag? In view of the Reynolds numbers involved ($Re \approx 10^6$) forget about Stokes with its linear dependence on speed v .

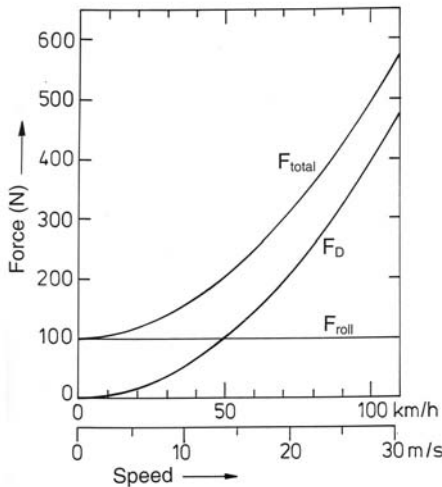


Image 4.1 | Rolling resistance, air resistance ('drag') and their sum, for a 1000 kg model car.

Instead, we should expect the drag F_D to be proportional to $\frac{1}{2} \rho v^2$, as already suggested by Bernoulli's law (ρ is the air density). On a vehicle with frontal area A , one can write $F_D = C_D \cdot A \cdot \frac{1}{2} \rho v^2$. Now, C_D is a complicated function of speed, but for the relevant v -range we may take C_D constant. For most cars, the value is between 0.3 and 0.4.

The total resistance is now shown in the figure, for a mid-size model car ($m=1000$ kg, $C_r = 0.01$, $C_D=0.4$ and $A=2$ m²).

It is funny to realize that the vertical scale immediately tells us the energy consumption. Since 1 N is also 1 J/m, we find that at 100 km/h this is approximately 500 kJ/km for this car. Assuming an engine efficiency of 20%, this corresponds to about 7 litres of gas per 100 km. At still higher speeds, the figure suggests a dramatic increase in the fuel consumption. Fortunately, it's not that bad, since the engine efficiency goes up, compensating part of the increase.

What about the engine power P ? Since $P = F \cdot v$, we find at 100 km/h about 15 kW. That's a moderate value. But note that, at high speed where drag is dominant, the power increases almost as v^3 ! Should we want to drive at 200 km/h, the engine would have to deliver 8-fold the power, or 120 kW. That's no longer moderate, I would say, and I'm sure the police will agree...

5

Old ears



If you are under, say, 35, you might as well stop reading: you should have no reason to worry about your ears. But for many of us who are somewhat older, a noticeable hearing loss may become a bit cumbersome every now and then. And as it turns out, the loss is worst where it hurts most: in the high frequency regime.

Let us first look at the data. In the figure, hearing loss data are given as a function of frequency for a large sample of people at various ages (Courtesy: Dr. Jan de Laat, Leiden University Medical Center). And indeed, already at age 60, the loss of high-frequency tones is frightening: over 35 dB at 8 kHz, increasing about 10 dB for every

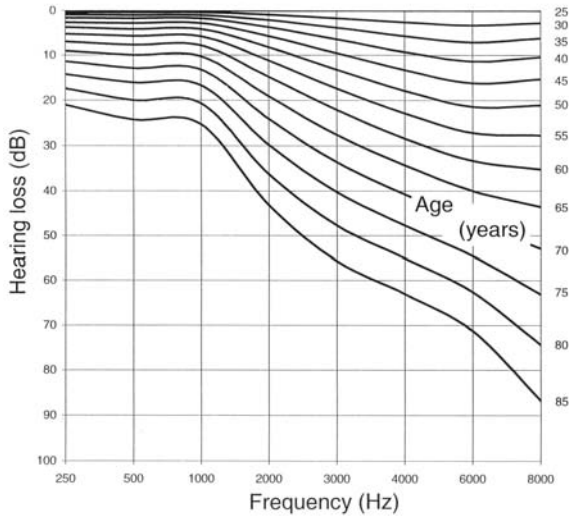


Figure 5.1 | Average hearing loss as a function of frequency, for persons aged 30 – 85.

5 years of age. Once we're 80, we'll be practically deaf for 8 kHz and up.

Why is hearing loss at the higher frequencies so bad? When listening to our stereo at home, we can turn up the treble a bit for compensation, no problem. And in a person to person conversation, we don't really have problems either, until we are having this conversation at some cocktail party. Then we notice: the background noise makes things worse.

One aspect playing a role here concerns consonants like p, t, k, f and s. They contain mainly high-frequency information, and will therefore easily be masked, or will get mixed up. Another aspect relates to the role of sound localization in selecting one conversation out of a background noise (sometimes referred to as the 'cocktail party effect'). We are pretty good at localizing sound: up to 1-2° in the forward direction (see William M. Hartmann in *Physics Today*, November 1999, p. 24 ff).

We use two mechanisms to do that. First, by using the phase- (or arrival time) difference between the two ears: the Interaural Time Difference (ITD). Of course, the information is unambiguous only if

the wave length is large compared to the distance between our ears. ITD is therefore effective only at the lower frequencies, say, below 1.5 kHz. However, in ordinary rooms and halls, reflected sound often dominates, especially for low frequencies. This is because the acoustical absorption decreases with decreasing frequency for almost all reflecting surfaces. As a result, the ITD becomes unreliable in such situations, and the low frequencies are not much of a help to spatially isolate one conversation from the noise.

Fortunately, we have a second mechanism, which uses the intensity difference between the two ears for sound coming from aside: the Interaural Level Difference (ILD). We remember that sound waves become effectively diffracted when their wavelength is much shorter than our head: the head casts a shadow, so to speak. Therefore, ILD works well above, say, 3 kHz.

Alas, look at the graph: the high-frequency region is where old ears have problems. So the ILD doesn't work too well either. In the end, we may have to resort to what deaf people do all along: use our eyes, and *see* the talking...

6

Fresh air



Whether at home or in the office: we feel comfortable when the temperature is around $20\text{ }^{\circ}\text{C}$ and the humidity around 50%. There's some interesting physics here, especially in wintertime, when we have to heat and – almost inevitably – to humidify the outside air. The humidity aspect is a trivial consequence of the steepness of the H_2O vapour pressure.

At $0\text{ }^{\circ}\text{C}$ and $20\text{ }^{\circ}\text{C}$, we find 6 and 23 mbar, respectively, almost a factor of 4 difference as seen in the figure. Therefore, when it freezes outside, the humidity cannot exceed some 25% inside, since the water content of the incoming air does not change by being heated. This is

so unless we add water to the room. The air-conditioning industry does that routinely in our labs and offices.

How hard is it to humidify the air in our home? In the stationary state this depends, of course, on the degree of ventilation. For a back-of-an-envelope calculation we use the rule of thumb that, for simple liquids including water, there is a factor of 1000 between the density of the liquid and that of the vapour if assumed at standard temperature and pressure. A litre of water, therefore, gives roughly 1 m^3 of vapour if it were at 1 bar (it gives 1.244 m^3 at STP, to be precise). Using the above 23 mbar at 20°C we find, for a room of 100 m^3 volume, that it takes about 1 litre to increase the humidity by 50% for a single load of air. If we assume a refreshment rate of once every hour, we see that humidification is effective only if we are prepared to pour a lot of water into our home daily, or we have to minimize ventilation.

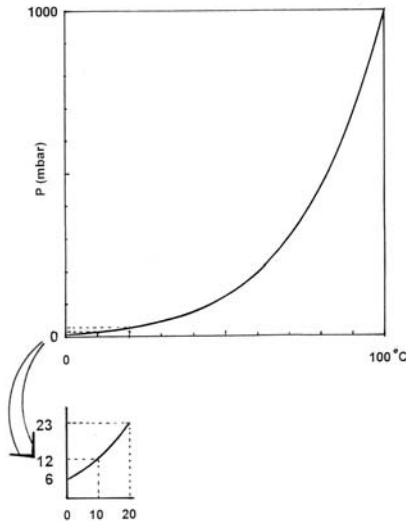


Figure 6.1 | Vapour pressure curve of water.

But ventilation is a must, if we don't want to run into health problems. In this context, an interesting physics aspect comes up. Suppose we instantaneously replace the air in our living room by cold outside air while keeping the heating off. Will the room be much colder after we wait for the new equilibrium to be reached? The answer is: very little, and it is easy to see why. It's all a matter of heat capacities, of course. But there is wooden furniture, brick walls, glass, metals etc. in the room, which seems to make an estimate pretty hopeless. However, if we're only interested in an approximate value, there is an easy way out. If specific heats are taken not per mass but per volume, values for most solids and liquids are pretty much alike (around $2\text{--}3 \text{ MJ.K}^{-1}.\text{m}^{-3}$). The reason is simple. We remember that atoms may differ enormously in mass, but they do not differ so much in 'size': the atomic number densities are rather equal in solids. Moreover, the contribution of each atom to the specific heat is roughly the same (around $3k$, with k Boltzmann's constant). For gases, of course, we have to take the above factor 1000 in the ratio of the densities into account.

Conclusion: when estimating heat capacities, a litre of liquid or solid and a m^3 of a gas at ambient temperature and pressure are pretty comparable.

So much for the rule of thumb. We can return now to our room. It is clear that the volume of the 'solid' content of the room is far larger than $1/1000$ of the air volume, even if we are honest and count only half of the wall thickness. This shows that, indeed, the temperature of the room will be hardly affected by a single load of fresh air. This trivial exercise also suggests that opening the refrigerator for a second or so puts about as much heat into the fridge as putting a tomato inside.

7

Diffraction-limited photographs



The optical performance of lenses, even in cheap cameras, is remarkably good these days. We don't have to worry too much about aberrations, even if we 'open up' and use the full lens aperture. Due to the steady progress in lens making over the years, our cameras – certainly the more expensive ones – are being gradually pushed to the diffraction-limited optics situation.

How does diffraction limit the resolution of our pictures? It all depends, of course, on the focal length of the lens (which we usually know) and the aperture, or effective lens diameter (which we may be unable to determine).

Fortunately, life turns out to be simple. Let us look at the textbook formula for diffraction through a circular aperture. When trying to image a point source on our film, we find that the radius of the resulting Airy disk is $1.22 \lambda f/D$, with λ the wavelength, f the focal length and D the aperture (the funny numerical factor 1.22 results from integration over rectangular strips).

The nice thing now is that the ratio f/D is the ‘F-stop’ value, which we recall having used on our non-automatic camera as one of the two parameters determining the exposure. The well-known series of values is 2; 2.8; 4; 5.6; 8; 11; 16; 22, spaced by $\sqrt{2}$, of course, in order to have double exposure between consecutive values.

Now, precisely *how* seriously are we limited by diffraction? Let us take a worst-case scenario, and assume that there is plenty of light such that the F-stop 22 is chosen. The formula for the Airy disk radius yields $r = 15 \mu\text{m}$ for the middle of the visible spectrum. In other words: we get a $30 \mu\text{m}$ diameter spot on the film, rather than a point. If we are using traditional, pre-digital-era 35 mm film, we may want to enlarge the 24 by 36 mm frame by a factor of 10 in order to have a nice size picture. This means that the diffraction spots become 0.3 mm in diameter, and are no longer negligibly small. The conclusion is that, if we use high-quality optics in our camera, it may be wise to open up the lens much further and use smaller F-stop values.

Now let us compare this to our digital camera: Is it the number of pixels that poses the limit to the resolution, or is it still diffraction? Using the above worst-case scenario with an Airy disk radius of $r = 15 \mu\text{m}$, and assuming the Rayleigh criterion for just-resolvable diffraction patterns (*i.e.*, a spacing by r is adequate to distinguish two adjacent ones from one another), we find that, on a 24 by 36 mm frame, we can store some 1600×2400 just-resolvable spots. If we were to image that pattern on our digital camera, and if we assume – somewhat arbitrarily – that the number of pixels on the chip must equal the number of the just-resolvable spots, we need almost 4 Megapixels. This is just about the performance of a standard digital camera. However, if we move from the $F = 22$ to the other extreme of $F = 2$, the diffraction-limited spot size shrinks

by a factor of 10. If the digital camera is to keep up with that, it has to increase its pixel number by a factor of 100.

So, when it comes to digital cameras, there is still room for improvement.

8

Time and money



Back in 1905, when Einstein was working on relativity in which 'time' plays such an important role, he would have never guessed that time would be measured with such an astonishing accuracy just a century later. As an example, think of GPS satellite clocks: to enable us to navigate with accuracies on the order of metres, their clocks have to be precise within nanoseconds. And in laboratories around the globe, laser-cooled Cesium and Rubidium fountain clocks reach an incredible fractional accuracy of about 6×10^{-16} . This translates into errors no larger than 20 ns in one year (which, coincidentally, contains almost exactly $\pi \times 10^7$ seconds).

But also in everyday life, things have changed dramatically. Most of us remember the pre-quartz era, when clocks rarely agreed to within a few minutes, and watches had to be adjusted every two days or so. Indeed, one had to resort to the radio if one wanted to know the exact time. By contrast, modern quartz clocks and watches routinely have accuracies better than 1 in 10^6 : some 30 seconds in a year. And, except for the switch-over to daylight saving time, adjustment is rarely necessary.

At what cost, in terms of kWh and Euros, do we read our daily time so accurately? The electrical energy consumption, even for a traditional analog clock operating on 230 V, is very small of course, as we can tell from the negligible amount of heat released. The electrical power for such a clock is typically on the order of 1W, and since a year has about 10^4 hours, it consumes about 10 kWh per year. In terms of money, that's about a Euro per year.

Now let us look at our digital watch. It typically operates on a silver oxide battery of 1.55 V having a charge of roughly 25 mAh. If we assume that the battery runs for at least two years, a back-of-the-envelope calculation shows that the watch operates on a power of less than 2 microwatt. That is very little indeed: it is six orders of magnitude less than an analog clock connected to the mains.

What about the cost? Such batteries cost, typically, 2 Euros, or a Euro per year of operation. Now lo and behold: isn't that what the analog counterpart in our home would cost?

The conclusion is simple. Our digital watches are very accurate and extremely efficient. However, the energy in their battery is extremely expensive, of the order of 50 000 Euros/kWh. But whatever type of clock we use for knowing the time as accurately as we do, the cost is 1 Euro at most for an entire year. If Einstein were alive today, he would probably agree: that's a lot of time for very little money.

9

Blue skies, blue seas

For the sky, it's simple. Most physicists know that the blue colour of the sky is due to the $1/\lambda^4$ dependence of Rayleigh scattering. But what about the blue of the sea? Could it be simply reflection of the blue skies by the water surface? That certainly cannot be the main story: even if the sky is cloudy, clear water from mountain lakes and seas can look distinctly blue. Moreover: those of us who like to dive and explore life under water will have noticed that, a few metres under the surface, bluish colours tend to dominate. Indeed, if we use an underwater camera and take pictures of those colourful fish, we notice that the nice red colours have almost

completely disappeared. And – unlike our eyes – cameras don't lie. We need a flash to bring out the beautiful colours of underwater life. In other words: absorption is the key: sunlight loses much of its reddish components if it has to travel through several metres of water. Or ice, for that matter: remember the bluish light from ice caves or tunnels in glaciers. And even the light scattered back from deep holes in fresh snow is primarily blue.

What causes the selective absorption of visible light by water? Spectroscopists know that the fundamental vibrational bands of H-atoms bound to a heavier atom, such as in H_2O , are typically around $3\text{ }\mu\text{m}$. This is way too long to play a role in the visible region. But wait: because of the large dipole moment of H_2O , overtone and combination bands also give an appreciable absorption. And they happen to cover part of the visible spectrum, up from about 550 nm , as seen in the figure.

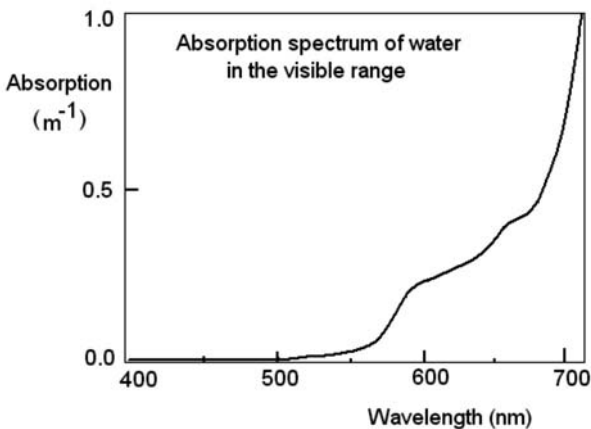


Figure 9.1 | Absorption of light by water.

The strong rise near 700 nm is due to a combination of symmetric and asymmetric stretch ($3\nu_1 + \nu_3$), slightly red shifted due to hydrogen

bonding (see, *e.g.*, C.L. Braun and S.N. Smirnov, *J. Chem. Edu.*, 1993, 70(8), 612). We notice that the absorption coefficient in the red is appreciable: it rises to about 1 m^{-1} around 700 nm, an attenuation by a factor of e at 1 m. It is no wonder that our underwater pictures turn out so bluish.

It is interesting to note: the spectrum of D_2O is red shifted by about a factor 1.4, since the larger mass of the deuterons makes for much slower vibrations. It is therefore shifted out of the visible region.

But that is not the whole story about the ‘deep blue sea’. For the water to look blue from above, we need backscattering. For shallow water, this may be from a sand bottom or from white rock. In this case the absorption length is twice the depth. For an infinitely deep ocean, however, we have to rely on scattering by the water itself and by possible contaminants. This may even enhance the blue color by Rayleigh scattering, as long as the contaminants are small compared to the wavelength.

If the water gets really dirty, things obviously become more complex. Scattering from green algae and other suspended matter may shift the spectrum towards green, or even brown.

But clear water is blue. Unless it’s *heavy* water, of course...

10

Cycling in the wind



When riding our bicycle, wind is bad news, usually. For one thing, it spoils our average speed when making a round trip. The reason is obvious: we spend more time cycling with headwind than with tailwind.

And what about a pure crosswind, blowing precisely at a right angle from where we are heading? That cannot possibly hurt the cyclist, one might think. Wrong. A crosswind gives rise to a much higher drag. Why is that? Don't we need a force in the direction of motion to do that?

So let us have a look at the relevant forces. The key is that air drag for cyclists is proportional to the relative air speed *squared* (just like for cars,

cf. page 20). This v^2 dependence spoils our intuitive feeling, as is easily seen from a vector diagram. See the figure, which illustrates the situation of a cyclist 'heading north'.

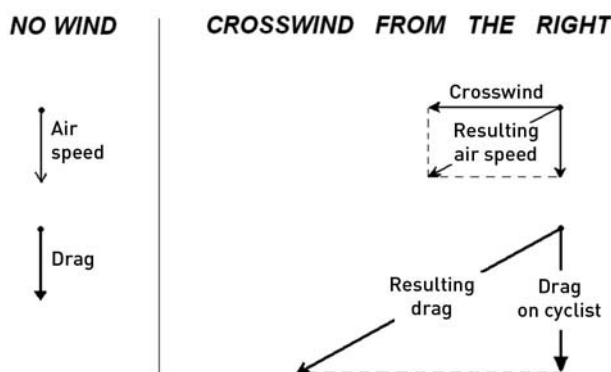


Figure 10.1 | Air speed and drag felt by a cyclist, in the absence of wind (left) and with a crosswind from the right.

The figure says it all. In the wind-free case (left) the cyclist feels an air speed equal to his own speed, and experiences a certain drag which we may call D . With a strong crosswind blowing from the East, the resulting relative air speed is much larger, and so is the drag. In our example, the resulting air speed is taken as twice the cyclist's speed (it comes at a 60° angle from the right). Consequently, the resulting drag is $4D$. So its component in the direction of motion is $2D$, or twice what it was in the wind-free case.

In order to *profit* from the wind, it has to blow slightly from behind. Of course, the angle for which the break-even point is reached, depends on the wind speed relative to that of the cyclist. In our example, where their ratio is $\sqrt{3}$, the break-even angle is 104.5 degrees, as calculated by Fokke Tuinstra from Delft University of Technology. But a pure 90 -degree crosswind always hurts the cyclist.

In fact it's even worse. Also the relevant frontal area, which determines the drag, is increased dramatically. It is no longer that of

a streamlined cyclist as seen from the front, but a $\sin \alpha$ projection of the cyclist plus his bike. And with α being 60° in the example, this is practically the full side view of the bicycle and his rider. Even the crouched position does not help much in this case.

Clearly, riding our bike in the storm is really brave. It makes good exercise. And it yields some funny physics, too.

11

Seeing under water



Most physicists realize that the human eye is not made for seeing under water. For one thing, if we open our eyes under water to see what's going on, our vision is blurred. The reason is obvious: since the index of refraction of the inner eye is practically that of water, we miss the refractive power of the strongly curved cornea surface. With its $1/f$ of about 40 diopters it forms an even stronger lens than the actual eye lens itself. Could we repair that with positive lenses? There is no need for a back-of-the-envelope calculation here: In view of the strong curvature of the cornea surface (radius about 8 mm), the idea of replacing it by a glass lens in a water environment is beyond hope.

We really need to restore the air-water interface in front of the cornea, and that is precisely what our diving mask does.

But there is more to it. Under water, our field of vision is reduced dramatically. Whereas we normally have an almost 180° field due to the refraction at the air-cornea interface, we lose that benefit once we're under water. The diving mask does not repair that since there is the compensating effect at the front of the mask. This is schematically indicated in the figure.

So, if you happen to be a scuba diver, beware! You have to turn your head much further than you may think necessary, if you want to be sure that you are not followed by a shark.

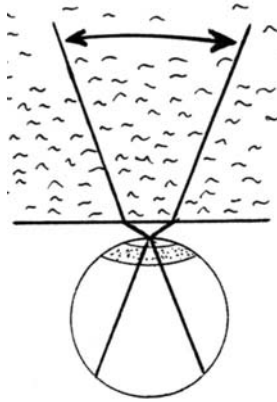


Figure 11.1 | Reduced field of vision under water for a person wearing a diving mask, schematically.

12

Cycling really fast



We remember that the cyclist on a horizontal road has to beat two forces (see p. 20). One is the rolling resistance, proportional to the total weight ($C_r mg$). The other one is air drag, proportional to the frontal area, the air density and the velocity squared ($C_D \cdot A \cdot \frac{1}{2} \rho v^2$). The two are equal at roughly 15 km/h for a normal bicycle. In view of the v^2 dependence, drag is by far dominant at record-breaking speeds. If you want to go fast, get rid of the drag.

One way to minimize drag is to use super-streamlined, recumbent bikes: HPV's, for Human Powered Vehicles. Their main advantage is a reduction of the drag coefficient C_D to 0.1 which is an order of

magnitude smaller than the value for a normal bike. As a result, speeds above 90 km/h have been a piece of cake for experienced riders ever since the 1980s. Indeed, in the U.S. during the nationwide speed limit of 55 mph (88 km/h), several riders earned an ‘honorary speeding ticket’ from the California Highway Patrol. More recently, in 1998, the landmark of 130 km/h was first reached by the Canadian Sam Whittingham.

For the real speed devil that’s not good enough. Why not abolish drag altogether, by riding behind a fast car having a large vertical board at its rear end (a technique also called *Motor Pacing*)? This is precisely what Dutchman Fred Rompelberg from Maastricht did in 1995, on the Bonneville Salt Flats in Utah, USA. He set off behind a powerful car on a special-design bicycle (but not an HPV) and reached a breathtaking 268 km/h. Sure enough, that made him the fastest man-on-a-bike ever.

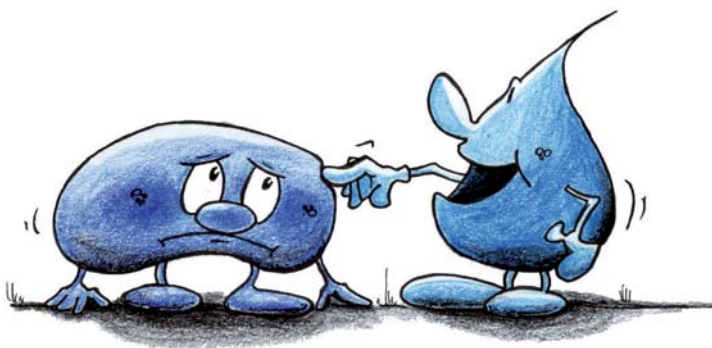
Now let us take this a bit further, by also reducing the rolling resistance. Let us do a thought-experiment and calculate how fast we could ride on the moon. Reasonable input data would be a peak power of 750 watt for the rider (which is what a trained cyclist briefly reaches on earth), a mass $m = 100$ kg (including the space suit), $C_r = 0.0045$ (a typical value for bicycles) and $g = 1.62$ m/s². Since the rolling resistance is the only force to be overcome, all we have to do is solve the equation $C_r mgv = 750$ W.

The resulting speed v turns out to be some 3700 km/h. That is really fast: over Mach 3 in terms of the terrestrial speed of sound at ambient temperature. But for lack of an atmosphere, we do not have to worry about sonic booms on the moon.

Much faster than that, however, may become a problem for prospective moon cyclists: 3700 km/h, that is about half the escape velocity...

13

Water from heaven



Large rain drops fall faster than small ones, that much is obvious for any physicist. But let's be a bit more precise. The terminal velocity follows from the balance between the weight of the drop and its air resistance. What exactly is the resistance of a drop falling through the atmosphere? We have to distinguish two regimes here.

If the droplets are real small, like cloud droplets (or fog particles, if you wish), the Reynolds number is so small that Stokes' formula applies: the air resistance is proportional to viscosity, radius and velocity: $F = 6\pi\eta Rv$. For a typical cloud droplet having a radius of 0.01 mm, we find a terminal velocity of about 1 cm/s. That is

very little indeed. But it goes up rapidly with size: since its weight is proportional to R^3 and the resistance only to R , the terminal velocity increases with the square of the size for such droplets. That applies to droplets up to about 0.1 mm in diameter, according to the handbooks (good old Ludwig Prandtl's book *Führung durch die Strömungslehre*, for example).

For ordinary raindrops, up from about 1 mm diameter, turbulent flow dominates. Here the weight is balanced by drag $F_D = C_D \pi R^2 \cdot \frac{1}{2} \rho v^2$, where πR^2 is the frontal area, C_D is the drag coefficient, which is about 0.5 for a sphere at the relevant Reynolds numbers, and ρ the density of air. For a rain drop of 1 mm in diameter, we find a terminal velocity of 16 km/h. Note that in this regime the velocity is proportional to the square root of the diameter. Consequently, a 3 mm raindrop reaches 28 km/h. And so on, we would guess. Given the above, we may expect that for the biggest drops – 5 mm, say – the terminal velocity is well above 35 km/h.

Wrong. Something interesting happens, as already noticed by German physicist Philipp Lenard a century ago. Using a vertical wind tunnel to balance the drop's speed, he noticed that drops larger than about 3 mm diameter become deformed like small pancakes, and have a flat bottom. Consequently, their frontal area is larger than for spherical droplets having the same mass. As a result of the increased drag, the terminal velocity hardly goes up any further: for raindrops of 4 and 5 mm it reaches an asymptotic value of about 29 km/h, practically the speed already reached by the 3 mm drops.

And beyond 5 mm? As soon as the diameter reaches about 5.5 mm, forces become so large that surface tension cannot hold the drop together, so it breaks up in pieces. So, for raindrops, there is no life beyond 5 mm.

14

Surviving the sauna



Human beings are not made for living in a 90 °C environment. And yet this is the temperature in the average sauna. How do we cope with such harsh conditions?

First, we use a towel to sit on, or we touch wood. Touching metal at that temperature is no fun at all. Even the glass door feels hot, although its thermal conductivity is far below that of metals, and its temperature is only about 60 °C, halfway between in and outside temperature. Second, the air is dry, which enhances cooling by perspiration. Incidentally, the dry air comes for free: due to the steepness of the water vapour pressure curve, even if the outside air

at, say, 20 °C is 100% humid, the humidity drops to 3% once the air is heated to 90 °C. If it freezes outside, that would even drop to 1%, provided that no water is added.

How fast would our body heat up, if we neglect perspiration? Let's do a back-of-the-envelope calculation. First, let's look at the conduction term. Since the effective air layer surrounding our body can be assumed to be around 3 mm thick, assuming a body surface area of 1.7 m² and a temperature difference of 50 K, we find some 700 W. Likewise, the radiation term yields about 800 W. So, conduction and radiation are roughly equally important, just like in normal circumstances. The difference is that they are reversed in the sauna. And they are an order of magnitude larger, due to the larger temperature difference and the fact that we have... eh... adapted our clothing. So the total heat load on our body is 1.5 kW, which corresponds to the power of an electric heater!

How fast will our body start to heat up? Taking a fair estimate for the heat capacity of our body of 200 kJ/K, we find a heating rate of 0.5 K per minute – as long as perspiration is negligible.

This is a sure way to disaster. So after a few minutes the sweating should begin. And it does, fortunately, even before we notice that our skin gets wet. Keeping up with the 1.5 kW heat load by sheer sweating would require 2.2 litre per hour. Our body will not be able to evaporate that much without forced air circulation.

Being physicists, we surely want to do a small experiment. Why not put some water on the stove, and see what happens? This makes us feel even hotter, and the question is why. Here is an educated guess. At least four contributions can be identified. First, the 100 °C steam coming off the stove is somewhat hotter than the sauna air. Second, it causes forced convection, which will heat our body even more, particularly if our skin is still dry. Third, the humidity goes up, which makes perspiration more difficult. And fourth: the thermal conductivity of water vapour is slightly higher than that of dry air, since water molecules are a bit lighter (and therefore faster) than N₂ or O₂. Which of those contributions is the

dominant one may be something to sort out next time we are in the sauna¹. After all, there is time enough to do some calculations and experiments.

But don't forget to keep an eye on the hourglass...

1. See also p. 59: The sauna – revisited.

15

Black vs. white

Does a dark-painted front door get hotter in the sun than a white-painted one?

“Of course”, says a layman, pointing out that a black surface absorbs solar radiation much better than a white one. “Should make no difference” says another layman with some science background, adding that a surface that absorbs well must also emit well. A physicist overhearing the conversation nods vaguely, remembering things like microscopic reversibility and detailed balance. But his intuition tells him it is not that simple. What precisely is happening here?

Let's do the experiment. On a bright and sunny day with little or no wind, we find the temperature of a white door to be 43 °C, and that of a dark green door 66 °C. The difference is clear: beyond the shadow of any doubt, the first layman was right.

Of course. It is obvious that the white paint remains cooler. The door surfaces absorb in the visible region of the EM spectrum, but emit in the infrared. According to Wien's law, the wavelengths differ – roughly speaking – by a factor of 20, *viz.*, the ratio between the temperature of the sun's surface and our ambient temperature of 300 K. That means that we are dealing with 0.5 μm for the incoming light *vs.* 10 μm for the IR emission.

The optical properties can vary dramatically over such a range. And they do. Almost all common surfaces are 'black' around 10 μm . If we look up their emissivity at such wavelengths, we find values near 1 for almost anything: common paints have values around or above 0.9, irrespective of colour. Even water and glass fall into that category, with emissivities well above 0.9. Metals, of course, are an exception. If they are clean and polished, such that multiple reflections are avoided, their emissivity is around or below 0.05.

But normal paint does not contain metal. The conclusion is, therefore, that the difference in temperature between the two doors is caused by their different *absorption* in the visible region. For the emission, all paints are equally black, except for Aluminium paints which can have emissivities below 0.3.

There is also a lesson here for our home heating. All radiators can be considered black, even the white ones: there is no need to deviate from our interior decoration taste as long as we stay away from Aluminium paint and the like.

What happened to the detailed balance argument? Obviously, detailed balance holds, but we have to consider one and the same wavelength. If we do, emission and absorption coefficients *are* equal. If, for example, copper looks reddish, it must absorb primarily green or blue. So, if we make copper *emit* visible light, *e.g.*, by introducing some copper salt into a hot flame, detailed balance tells us that it should emit green or blue. And sure enough, it does.

16

Hearing the curtains



If there is one place in the house that we like to sing in, it's the bath. The reason is the unusually long reverberation time: the exponential decay of any sound is slow. It is all described by Sabine's law, which states that the typical decay time of sound in a room is proportional to the volume of the room and inversely proportional to the total area of surfaces that completely absorb the sound (like an open window would, for example).

So, if it comes to reverberation, the bathroom is pretty unique. It usually has bare walls, tiles on the floor, and little or no furniture that could absorb sound. Even our own clothing may give a much

smaller contribution than usual – if it gives any contribution at all.

That leaves only the curtains as an efficient means for absorbing the sound, if we assume that we do have curtains in our bathroom. So if we *really* want to enjoy our own singing, we would probably be well advised *not* to close the curtains but leave them open, in order to have as small an absorbing surface as possible and hence a maximum reverberation time. That seems very plausible, if we follow our physical intuition.

Wrong. Whether we have the curtains open or closed makes very little difference for the sound absorption, and hence for the reverberation. The reason is somewhat subtle: the sound is indeed dissipated at surfaces, by friction losses of the sound waves near the surface. But, more precisely, we have to consider the microscopic surface of the material, which includes the pores. That is the reason why porous media like thick draperies, carpets, fibrous mineral wool, glass-fibre and open-cell foam are usually good sound absorbers. And for the curtains this means that, as long as the sound waves have easy access to the inner surface, it does not matter much whether the curtain is spread out over the entire wall or bundled together in a corner of the room.

The conclusion therefore must be: with our eyes closed, we can't really tell whether the curtains are open or closed. We do notice, however, if they happen to be at the dry cleaner's.

17

Fun with the setting sun



The setting sun plays a few tricks that any physicist will appreciate. One of these is well known: the sun appears unusually red when setting. It is the $1/\lambda^4$ dependence of Rayleigh scattering which selectively removes the blue end of the spectrum from the transmitted light.

Less well-known is the fact that the sun is not where it seems to be, during sunset. In fact, it may be behind the horizon while we still see it. We are not talking here about the finite speed of light, which makes us see the sun about 8 minutes late. We are talking about the refraction of the sunlight due to the vertical gradient in the index

of refraction, which in turn is caused by the density gradient. If we ignore temperature gradients for a second, the density decreases with height due to the decreasing atmospheric pressure, by a little over 1% for every 100 metre, *i.e.*, $n^{-1}(dn/dz) \approx 1 \times 10^{-4}/\text{m}$. As a result, the light rays are bent downward, in the direction of the earth's curvature. This may be seen as the inverse of the well-known 'highway mirage', the apparent pools lying across the pavement when the sun shines.

Granted, temperature effects can be much larger than the barometric pressure effects, which is easily seen if we realize that, for constant pressure, we have $n^{-1}(dn/dz) = -T^{-1}(dT/dz)$. But let us look at what happens if temperature gradients are negligible, or – even nicer – if temperature *increases* with altitude. Then the temperature effect – if any – adds to the barometric pressure effect. Now the light rays tend to follow the earth's curvature, which makes us see the sun just after sunset. This effect occurs both at sunrise and at sunset, and adds an extra 5 minutes of daylight to each day. Note that the finite-speed effect mentioned above does not do that; it just gives an 8-minute offset throughout the day.

Since bending of light rays in the atmosphere is stronger for lower-lying rays, there is a second phenomenon: the sun appears to be flattened by about 10%. The fact that we do not always notice this is due to the competing effect of temperature.

Finally, there is the somewhat mysterious 'green flash' that people sometimes observe at the moment of sunset. It lasts only for a few seconds, and requires somewhat favourable atmospheric conditions. Why green, and why only for a few seconds? There are a few things here that we have to combine. First the refraction, which makes us see sunlight after the actual sunset. Due to dispersion this effect is strongest for the blue end of the visible spectrum. This means that we expect blue to be visible longest, while the red end of the spectrum has long disappeared. But blue light is almost absent in the setting sun, as seen above. The result is that the last flash of sunlight is dominated by green.

The green flash: a last good-bye from the setting sun. But at least a good-bye in style.

18

NOT seeing the light



Back in 1808 the young French soldier Étienne Louis Malus noted that there is something funny about reflections. While looking through a crystal of Iceland spar calcite in his Paris apartment, he noticed variations in the sunlight reflected from windows in the *Palais du Luxembourg* across the street when he rotated the crystal. This observation, often considered as the discovery of light polarization, laid the basis for our Polaroid glasses. Indeed, the most common use of Polaroid glasses is aimed at reducing annoying reflections.

To fully appreciate the issue, let us recall the behaviour of light when reflected from glass, or from water. The reflectance as a function of

incident angle θ (the angle to the normal) is given here for convenience. The graph is for the case of glass, but is only marginally different in the water case. It shows the reflectance for the two polarizations parallel and perpendicular with respect to the plane of incidence. The dashed curve is the average, or the effective reflectance for non-polarized light.

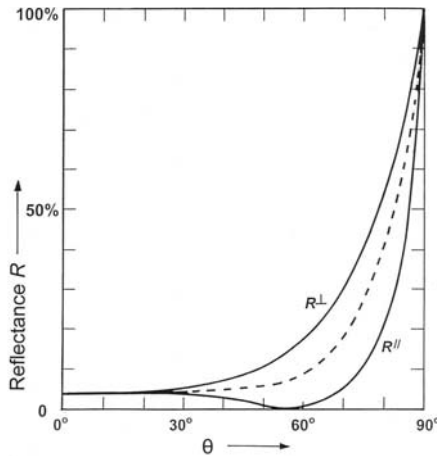


Image 18.1 | Reflectance of light from a glass surface vs. incident angle for two polarizations and for unpolarized light (dashed curve).

Before entering into a discussion of the two different polarizations, it is interesting to notice that for grazing incidence ($\theta = 90^\circ$) the reflectance becomes unity. Therefore, the image of the setting sun above a quiet lake appears just as bright as the sun itself, for example.

At the other end of the axis, for light incident along the surface normal, the reflectance is a few percent only: for glass having a refractive index $n = 3/2$ we find $(n-1)^2 / (n+1)^2 = (1/5)^2$ or 4%. For water with $n = 4/3$ we find even less: $(1/7)^2$ or 2% only. Therefore, if we look straight into a pond, the reflection of our own face is really weak, and there is a fair chance that we can see the fish, provided that they are there and that the water is clear.

But we can do better than that by going to angles in between these two extremes and using Polaroid glasses. Obviously, our best

choice is Brewster's angle, where one of the two polarizations has zero reflectance, such that the reflected light is completely polarized. It is the angle whose tangent is the index of refraction: $\theta = 56^\circ$ for glass and 53° for water. Here our Polaroid glasses work perfectly.

So, if we want to make a picture of something behind glass, Brewster comes to our rescue, provided that we orient our Polaroid filter correctly. And, in the case of the pond: using Polaroid glasses we can completely get rid of the reflection of the sky. Use a bit of physics, and outsmart the fish.

19

Thirsty passengers



As a rule of thumb, commercial aircraft consume some 10 cm^3 of fuel per seat per second. That sounds like a lot. Imagine the whole cabin taking a sip each second, with the flight attendant beating time. Funny. But that's what the fuel consumption amounts to.

No wonder, one might think: at such tremendous speeds the drag must be enormous. Compare that with the slow boats of yesteryear, which took a week to cross the Atlantic. They must have been a lot less wasteful than those fast planes nowadays.

But wait: shouldn't we look at fuel consumption per *kilometer* rather than per second? Back to the rule of thumb: 10 cm^3 per second

is 36 litres an hour, during which the plane flies some 900 km. That yields 4 litres per 100 km. Modern efficient aircraft do a bit better than the rule of thumb, and arrive at, say, 3 litres per 100 km per seat. So, two passengers consuming a joint 6 litres per 100 km are just as wasteful as if they were sharing a reasonably efficient car.

What about the slow boat? Surprise. A large passenger boat or a cruise ship consumes about 25 litres per 100 km per passenger. Despite its moderate speed, the boat is much worse than the plane, in terms of fuel consumption per passenger km. How come? A bit of physics leads the way. Of course, the drag is determined not only by speed but also by the density of the fluid. Water and air differ by three orders of magnitude. It's even more than that: since commercial aircraft cruise at 10 km, and since the density goes roughly as $\exp(-h/8\text{km})$, they cruise at roughly $1/4$ of the standard value.

But perhaps the biggest difference is the payload. On a cruise ship, the mass of the passengers plus their luggage typically amounts to a few tenths of a percent only. The reason, of course, is that a cruise ship is a floating village, with shops, restaurants, swimming pools and the like. Even a huge modern vessel like the Queen Mary 2 with its 150 000 tons carries 2600 passengers only. Compare that with a big airliner. The total mass of its passengers is well above 10% of the aircraft.

Agree: in the interest of energy and the environment, we travel way too much. Kerosene is way too cheap, and we fly way too often. But if we *have* to, crossing the Atlantic by boat would be even worse.

20

The sauna – revisited



On page 45 we addressed several aspects of the funny situation in which we find ourselves when visiting the sauna. One question remained a bit open: What exactly causes the temporary heat pulse that we feel when we pour some water on the hot stones, thereby temporarily elevating the air humidity? All we could produce at the time was the ‘educated guess’ that at least four different contributions could play a role. But that did not really solve the question satisfactorily. Fortunately, Timo Vesala, professor of Meteorology at the University of Helsinki, came to the rescue. Having done qualitative observations during a few years in his

own sauna twice a week, he solved this non-trivial problem and published a paper on the issue²).

Here is the surprise: The latent heat released in the condensation of water vapour onto the skin is an important mechanism – perhaps the most important one. The reason is that our skin is most probably the coldest place in the sauna, and the humidity can easily become 100% near the skin.

That's beautiful! We are used to think in terms of evaporation from our skin, not condensation. But the sauna is something special, and we should think beyond the box.

To check the validity of the argument, let us assume the sauna temperature to be 100 °C (real Finnish sauna's are somewhere between 80 and 110 °C). This eases the analysis since 100% humidity

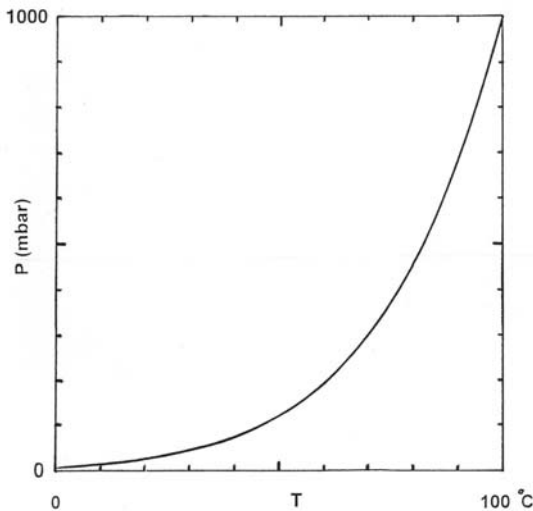


Image 20.1 | Vapour pressure curve of water.

2. T. Vesala: Phase transitions in Finnish sauna. In: *Nucleation and Atmospheric Aerosols*. Proceedings of the Fourteenth International Conference on Nucleation and Atmospheric Aerosols, Eds. M. Kulmala and P.E. Wagner, 403-406. Pergamon, 1996.

nicely corresponds to a vapour pressure of 1000 mbar. What happens can now be easily seen from the vapour pressure curve, reproduced here for convenience.

As pointed out in the previous sauna item (and readily checked from the curve), the humidity will automatically be 3% at most if the outside air is heated to sauna temperature. This will go up if extra H_2O is released, for example by perspiration. In practice, the average humidity seems to be 8% or so).

From the water vapour pressure curve we see that 8% humidity (*i.e.* 80 mbar in this case) will lead to saturation at about 40 °C. This is almost exactly the sauna-warmed temperature of our skin: 43 °C), as confirmed by infrared skin-temperature observations. In other words: if the humidity is raised a bit, to 10% for example, condensation onto our skin will be inevitable. And this is precisely what happens when we throw water on the hot stones.

Timo Vesala has also done a quantitative analysis to assess the extent to which condensation contributes to the heat pulse. He finds that this latent heat flux is around 4 kW. This is on the same order of magnitude as the ordinary heat flux, which is already enhanced during the heat pulse by the extra convection (the analysis does not include heat exchange by radiation between the body and the hot oven).

So, next time you're in the sauna, you may want to check the perspiration experiment.

But if you prefer to just sit and relax: no sweat.

21

Refueling



We don't usually think about it when driving down the highway, but what will traffic look like after the fossil-fuel age? How will our great-grand-children move 'in the fast lane'? No longer in a gasoline-powered car, probably. An all-electric car perhaps? Or a hydrogen car powered by fuel cells? Or will they use some synthetic liquid fuel to power their engine? Things don't seem very clear yet.

Let us assume for a moment that it will be an all-electric car. Sure, there is a problem with the weight of the batteries: even with the best battery type now available, the weight of our car would roughly double if we want to carry batteries with the equivalent of 50 or

so litres of gasoline. But let us be optimistic: let us assume that we are able to improve the energy density of batteries by another order of magnitude. That would make the extra weight quite acceptable. Problem solved, one would think.

But now another interesting aspect comes up. How about refueling? When driving long stretches on vacation in our present cars, refueling is a piece of cake. We can do it during the coffee break, for example. Now let us consider the electric car. Suppose our batteries are running low, and it is late afternoon. Fortunately, our hotel is near. No need for a gasoline station: there are power outlets in the hotel, and we will nicely reimburse the hotel owner. But how long will the charging procedure last, if we want to drive another 700 km the next day?

Let us do a back-of-an-envelope calculation. A standard power outlet can draw 16 A at most if we don't want to blow the fuse. At 220/230 V this yields a power of, say, 3.5 kW.

Compare this with the average car driving on the highway: it uses about 15 kW, which is higher by a factor of 4. The conclusion is that we need to charge the batteries for roughly 4 hours for every hour of driving. Since we want to drive for about 7 hours the next day, we need to charge for about 28 hours.

So if our great-grand-child will be driving an electric car, he had better pick a hotel that is especially equipped for fast overnight charging. And become very good friends with the hotel owner. Otherwise: forget about an early start the next morning.

Given the above result, it is interesting to calculate the energy flow into our present car when we fill our tank with gasoline. It turns out that we pump about 0.6 litres of gasoline each second. With the heat of combustion being about 35 MJ per litre, this translates into $21 \text{ MJ/s} = 21 \text{ MW}$ (!). In terms of electric power, given a conversion efficiency of $\frac{1}{3}$, this is some 7 MW. That is 2000 times as fast as charging batteries from a standard electric outlet, see above.

Should our grand-grand-children muse about such numbers when driving down the highway, chances are that they'll look back at us and our petroleum age, and think: gee, weren't *those* guys lucky...

22

Counting flames



Ask any layman this simple question: “If you run your hot-water tap, you are using energy, right? How many lights do you think you could switch on from that energy during the same time?”

The answer will probably be something like: ‘Well, let’s see, I guess 10, or perhaps even 20’. He or she will be surprised if we say that it may be as many as 1000.

The layman does not know that the specific heat capacity of water is remarkably high. And he or she does not realize the full extent of the first law of thermodynamics.

For us physicists, it's easy. We could even explain things by counting flames, knowing that a small flame produces about 100 watts. Take a match. Its mass is about 0.1 grams and therefore its wood contains roughly 2 kJ. Now just assume that it burns for about 20 seconds, and there you go: $2000 \text{ J} / 20 \text{ s} = 100 \text{ W}$. For a candle we can do the same exercise. Find out for how long it will burn, look up the heat of combustion of paraffin or stearine, and again: about 100 W. So the rule of thumb is simple: a small flame is a heater of about 100 watts.

From here it's downhill. First let us look at a camping gas cooker, or – if we use natural gas at home – at the gas stove. Each burner has 20 to 30 flames, so a burner should produce 2 to 3 kW of heat. And sure enough: if we look it up on the internet, Google tells us that our guess was right. For a hot water tap, though, that is not enough. If we happen to be familiar with gas geysers, we remember that they have about 10 rows of 10 flames. That makes $100 \times 100 \text{ W}$ or 10 kW. And the gas geyser isn't even a device of great luxury. It has too small a capacity to produce a decent shower, for example. It is therefore safe to assume that an average hot water tap will easily exceed those 10 kW.

But let us not overdo things, and stick to 10 kW. And let us for simplicity assume that the water is heated electrically, so we can directly compare that to electric lighting. An efficient light source producing 600 lumens consumes about 10 W. That is, indeed, a factor of 1000 lower than the hot-water tap.

This is a nice little lesson for any layman concerned about energy and global warming. Most often, he or she identifies energy use at home with things that turn or things that shine: electric motors or lighting.

Wrong. It's not motion. It's not light. *Heat* is our guide.

23

Drink or drive



For most of us, the small revolution went unnoticed. When we drive our car these days, our gasoline engine is no longer running on fossil fuel, *i.e.*, on the solar energy harvest of millions of years ago. For 10% or so, it is running on the solar energy harvest of *last* year: on bio-ethanol, that is. For diesel engines, it may be even more than just a few percent. In Germany, for example, up to 200 000 cars have been running on pure biodiesel lately. And the European Commission's goal is that 10% of all transport fuel should be biofuel by the year 2020.

The EC may have been ill advised to set that goal. We have witnessed a dramatic increase in food prices world wide over the last years, and

part of that is due to biofuels. Is the whole idea of biofuels just a hype, then?

Let us do a back-of-an-envelope calculation. Our daily food amounts to about 10 000 kJ/day. In terms of oil or gasoline, that is $\frac{1}{4}$ litre per day, only a small fraction of what our cars needs as a daily diet. In other words: If adopted on a world wide scale, this idea is bound to run into problems, if we assume that the biofuel competes with food, which is what the present, so-called first generation, does.

If we utilize not only the food-related part of plants, but the whole harvest of photosynthesis including straw and the like, we can do better. But even this 'second generation' of biofuels has its limits. The basic reason is that the overall efficiency of natural photosynthesis is low. In a typical European climate, it is somewhat below 1% as an average over the yearly solar energy influx. Which is bad news for densely populated and energy-intensive countries. Take the Netherlands, for example. Even with an optimistic photosynthesis efficiency of 1%, the area needed for the total energy consumption to be based on biomass on a sustainable basis would be more than twice the total area of the country.

Granted, sooner or later we will have to rely on the sun for an appreciable part of our energy supply. But can't we do better than good old photosynthesis? Think of photovoltaic cells, for example. Crystalline silicon cells routinely have an overall efficiency of 10 to 15%, while multi-junction concentrator cells achieved a record 43% in the summer of 2007. This suggests that we may be better off relying on high-tech solutions, rather than trying to meet the energy demand of the modern energy-intensive society with methods of the Middle Ages.

In any case: Should we base our future fuel consumption on bio-ethanol, we sure would run into nasty dilemmas. For example, during the reception of the 50th anniversary of the EPS in 2018, we would face questions like 'Shall we have another glass, or shall we drive our car for another 300 metres?'

24

Feeling hot, feeling cold



Even on a cold day, a bit of sunshine can make a tremendous difference. People will say things like ‘It is supposed to be 15 °C according to the forecast, but in the sun it’s at least 25’. Although this may contain some truth in terms of heat balance, it is, strictly speaking, nonsense. There is no such thing as ‘temperature in the sun’. How would one measure that? Different types of thermometers hanging in the sun would give widely different readings, depending on construction, optical properties and the like. The only decent definition of air temperature is derived from the mean kinetic energy of the molecules: $\frac{1}{2} m \langle v^2 \rangle = \frac{3}{2} kT$. Radiation has nothing to do with it.

But measuring the kinetic energy of the molecules in a gas directly is not exactly a piece of cake. Therefore we use an indirect way: the thermometer. It's easy to use, but not always reliable. The problem is the low thermal conductivity of air. This makes the thermal contact between the air and the thermometer very poor. As a consequence, the influence of radiation is hard to suppress. If the thermometer is in the sun, forget a reliable measurement. But even in the shade, indirect radiation will cause our thermometers to be slightly optimistic. No wonder that meteorologists have strict rules for determining the temperature: thermometers must be placed inside well-ventilated casings, which are painted white, placed 1.5 metre above the ground, etcetera. If you think about it, it's almost a miracle that air temperatures are accurately measured at all.

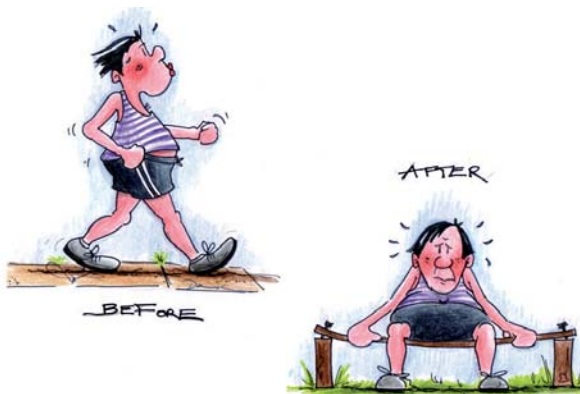
Wind is another source of misunderstanding, if it comes to temperature. Obviously, if the wind blows around our body (or, in fact, around *any* object that is heated above ambient temperature), the heat losses by conduction will increase. The reason is that the insulating layer of air – normally a few mm thick – will become thinner once the wind blows. The effect is the same as if the air temperature were lower. That seemingly lower temperature is often called the 'wind chill' factor. Although this is a widely known concept, many people are still missing the point. An example is the journalist who concluded, using the wind chill table, that the water in his car's radiator would freeze well above the freezing point, if only the wind would blow...

If we think about it, wind chill is an ill-defined concept. For one thing, it depends upon the clothing that we wear. For example, in the limit of infinite insulation, wind would not bother us at all, and the wind chill factor would become meaningless. All we can say for sure is that any correction for wind must asymptotically reach a limiting value if the wind speed goes to infinity. Consider bare skin: eventually, our skin would assume the air temperature, and the heat losses would be limited only by conduction inside our own body. Not an appealing prospect, if it freezes outside.

Sun and wind both make the concept of temperature a bit fuzzy. Thank heaven that kinetic theory provides us physicists with a reliable definition, come rain or shine.

25

The way we walk



Centuries of evolution have given mankind plenty of time to learn how to walk. Walking is a reasonably efficient way of getting around, although not nearly as efficient as riding a bicycle. A few obvious features help us to walk efficiently: we move our arms and legs in antiphase, thus keeping the total angular momentum more or less zero. And we swing our legs at almost the natural pendulum frequency, which is around 1 Hz for adults. Indeed, traditional military marches proceed at 120 steps per minute: exactly 1 Hz. Given a standard step length of 83 cm, the corresponding marching speed is almost exactly 100 m per minute. Beautiful! This result does not

serve to illustrate the superiority of the metric system, but it sure is handy to know when hiking.

Energetically speaking, walking on a horizontal surface is a special case. We have no external force to overcome, in contrast to climbing the stairs, for example, where we have to fight gravity to increase potential energy; or to rowing and cycling, where we have to overcome drag from water or air. Walking is different: even aerodynamic drag is negligible (remember that it is proportional to the square of the speed). All energy that we produce is dissipated by our own body.

One may wonder why walking costs any energy at all. In fact, experiments show that the metabolic cost of walking, derived from oxygen consumption and carbon dioxide production, is about 2.5 W per kg of body mass. This is roughly 200 W for an adult. Why is that still so much? It is because human walking is mechanically complex. It involves the activity of numerous muscles, and various theories are being developed to arrive at a comprehensive description.

As innocent physicists we may offer an obvious clue: the effective displacement may be horizontal, but our centre of mass must be raised by some 4 cm at each step. Could that account for the high metabolic cost? This simple explanation, tempting though it may be, is not supported by the evidence. Experiments by Arthur D. Kuo at the University of Michigan have shown that a walking pattern that reduces the displacement of the centre of mass, does *not* reduce metabolic cost. In fact, it makes the metabolic cost go up. Also when the step length is varied beyond our natural step length, the cost goes up. In other words: the way we normally walk is also the most efficient one.

The conclusion seems to be inevitable. If we really want to walk more efficiently, we should not try to improve on our steps by thinking physics. We shouldn't even think at all, just walk. And if we are still not satisfied with the result, there is only one alternative: go home and pick up our bicycle...

26

Wine temperature



It happens to all of us, once in a while: We go to open up a bottle of wine only to discover that the temperature isn't right. If it is a red wine and it's too cold, it's easy: we can simply put it in the microwave oven for a few seconds (don't tell the wine maker). But if it's a white wine that is much too warm, we have a problem. All we can do is put the bottle in the refrigerator – and be patient. How long will it take for the wine to reach the desired temperature?

Being physicists, we realize that the answer is determined by exponential decay of the temperature difference between bottle and refrigerator, with the time constant being the thermal relaxation time of a bottle of wine in

air. This is not exactly a problem treated in the textbook, but it is easily solved. The thermal relaxation time simply equals RC , the product of thermal resistance R and heat capacity C . Since C of the wine can be well approximated by the value of water, all we need to find out is the thermal resistance R of the glass layer between wine and air. This may sound cumbersome for an exotic shape like a wine bottle, but an approximation in terms of a parallel-plate geometry will do. Therefore we put $R = d/kA$, where d is the glass thickness, k its thermal conductivity coefficient and A its total outer surface area.

The calculation is easy to do, with d just over 3 mm as determined from the weight of an empty bottle, its outer surface area and the mass density of glass. The resulting relaxation time is found to be almost exactly 3 minutes.

Three minutes! Or even less if we include radiation! That can't be right, as we know from experience. And indeed, it isn't. The reason is that we have grossly underestimated the value of R . We must include the thin layer of air surrounding the bottle, which is effectively convection-free, and in which thermal transport relies on conduction. This layer represents a much larger thermal resistance than the glass does. The 3 minutes just calculated must be considered a lower limit. The experimental value of the relaxation time of a wine bottle in the refrigerator is found to be about 3 *hours*, corresponding to an effective air layer of a few mm thick.

Of course we could speed up the cooling process by putting the bottle in iced water rather than air. If natural convection in water *and* wine is sufficiently effective, we may approach the limit calculated above.

But there is a better way to cool our wine: use a commercial 'cooling jacket' which contains a cooling gel that has a large latent heat capacity. The advantage is that it may be pre-cooled to far below 0 °C *and* that it provides a good thermal contact with the wine bottle.

Obviously, this trick works for *any* bottle, regardless of its contents. But in view of the heavy brain work we have just done, it seems only fair to treat ourselves to a delicious Pouilly-Fumé from the Loire, or a Pinot blanc from Alsace. At *just* the right temperature... after only 8 minutes.

27

Over the rainbow



Everybody knows the rainbow, and most physicists know its optical background.

But there is one question about rainbows that even most physicists cannot answer off-hand: what about the brightness of the sky above and below the rainbow?

In order to find the answer, let us first remember how the rainbow itself comes about. Geometrical optics will do, if we assume the size of rain droplets to be large compared to the wavelength of light. The key is that light rays making one internal reflection inside a raindrop have an extreme in their deviation as a function of ‘impact parameter’

if we put it in molecular collision language. That is, outgoing ray no. 2 in the figure makes the largest angle with respect to the horizontal, although it has incoming neighbours at either side (a fact that can easily be demonstrated by slowly moving a cylindrical glass of water through a laser beam). Consequently, when we turn away from the sun and look at a rain cloud illuminated by the sun, the reflected light is extra bright at this angle of about 42 degrees with respect to the sun's rays: the *rainbow angle*. Due to dispersion, the angle is different for each colour, and we see a colourful cone of light: the rainbow.

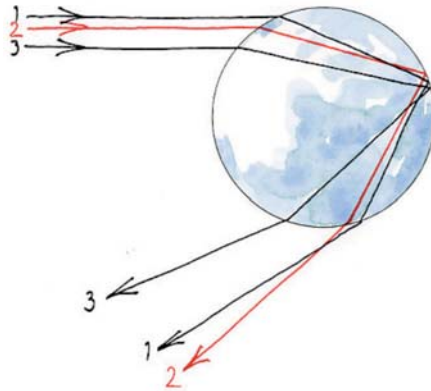


Image 27.1 | Sunrays leaving a raindrop after a single internal reflection, responsible for the main rainbow. Ray 2 illustrates the extreme in the deviation, at the 'rainbow angle' of 42° .

So far for the rainbow itself. Now what about the brightness of the sky next to it? From the figure it is obvious that there is also light reflected at angles *smaller* than 42 degrees, but not at larger angles. Conclusion: the sky is brighter inside than outside the rainbow.

But wait: this was only about the primary rainbow. What if there is also a secondary rainbow, having an angular radius of about 52 degrees? We recall that the secondary bow is caused by the extreme in the deviation of rays which leave the droplets after two internal reflections. It has inverted colours since the light rays have turned the other way around inside the droplets.

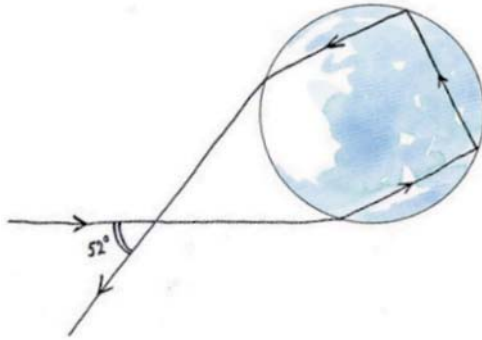


Image 27.2 | Sunrays leaving after two internal reflections, causing the weak secondary rainbow at a 52° angle.

How does the secondary rainbow affect the brightness of the sky? Interesting question, but easy to answer if we start by looking at rays going through the center of a droplet (impact parameter zero, or a ‘head-on collision’). After two internal reflections, such a ray continues to move along its original trajectory. With increasing impact parameter, the outgoing rays will gradually move over toward the incoming direction until they reach their extreme: the (secondary) rainbow angle of 52 degrees. Consequently, they will not reach the ‘dark’ area in between the two rainbows. So the conclusion emerges that the sky is brightest inside the primary and outside the secondary bow. Complicated though it may seem, it reminds us of a well-known song. The sky is bright, *somewhere over the rainbow*. But not *everywhere*.

28

New light



The old-fashioned incandescent light bulb with its tungsten filament is a marvelous piece of technology. If we switch it on, it needs only a split second to light up our office, our home or our fridge. Sure, this instant reaction is largely due to the low heat capacity of the filament. But there is more to it than most of us realize: the resistance of tungsten, like all ordinary metals, has a positive temperature coefficient. Indeed, if we calculate the resistance from the bulb's power and the grid voltage, and compare it with a direct measurement at ambient temperature, we find that the hot filament has a larger resistance by a factor of 20 or so. This

means that, if we switch the bulb on, the initial power is very high, making the bulb rush to its operating temperature in no time at all. And the other nice thing is: should the voltage go up for some reason (which it did, by the way, from 220 to 230 V over the past few decades) the voltage surge will be counteracted by the increased resistance. This dampens the power increase, and allows the bulb to withstand the surge. Bulbs from the good old '70s or '80s should have no problem adapting to the 21st century.

Alas! The efficiency of the incandescent light bulb is downright lousy. It is so poor, that the members of the European Parliament recently decided to ban the bulb. They have a point. There is no way we can ever make a glowing piece of tungsten into an efficient light source. For one thing, the emission peak at 3000 K is around $1\text{ }\mu\text{m}$ wavelength, as follows directly from Wien's law. The corresponding emission curve has only a small overlap with our eye's narrow sensitivity curve at around $0.5\text{ }\mu\text{m}$. And if we go much higher than 3000 K, the filament won't last very long. By invoking halogen vapour to redeposit evaporated tungsten back onto the filament, we may get a bit closer to the melting point of 3700 K. But even if we were able to find a high-melting-point metal which could be heated to 6000 K (roughly the effective solar temperature, with an emission peak that nicely fits our eye sensitivity), its black-body radiation curve would still be much broader than the eye sensitivity curve, with a lot of energy wasted.

What we need is a smart light source which selectively emits radiation that our eyes can see. And which has no filament that slowly but surely evaporates.

So we turned to gas discharge and invented fluorescent TL lighting long ago, with an efficiency of 100 lumen per watt, and – more recently – its folded version known as the energy saving lamp, reaching 50 lm/W. And, of course, the Light Emitting Diode as its solid-state counterpart, with a similar efficiency, depending on the type. Compare this to a poor 12 lm/W for the good old incandescent bulb!

We may not always be happy with politics: when it comes to lighting, however, we have to admit that 'Brussels' has a point.

Incandescent light bulbs may be fast and convenient, their emission spectrum may be nice and continuous, but in terms of efficiency, they are beyond hope. It's about time to kiss those bulbs goodbye.

29

Windmill nuisance



Many people dislike them, and some find them downright awful: it's the wind turbines scattered all across Europe these days. And let's be frank: there is nothing quite like the relentless droning of rotor blades to spoil the peace and tranquility of the countryside. Why do we put those things all over the place? As physicists we realize that wind power is proportional to v^3 , with v the wind speed. So it may not be such a great idea to put those turbines on land, let alone in the middle of continental Europe where wind speeds are typically low. Why not put them off-shore where winds are strong, in the North Sea or the Baltic, for example? A few off-shore wind

farms have already been put into operation recently, and a number of others are planned. Shouldn't we forget those monsters on shore altogether?

Let us have a closer look at the two options. First: wind turbines at sea. How many do we need, to begin with? Let us assume we want to have the equivalent of, say, 1500 MW, which is typically the electricity output of a large conventional or nuclear power plant. Modern wind turbines with a rotor diameter of 90 metres can produce 3 MW each. One may now be tempted to conclude that we need 500 turbines. Wrong. We have to include the load factor, *i.e.*, the average output divided by the maximum output. This is typically 30 to 33% for wind turbines at sea (and up to 25% on shore). So we need about 1500 turbines of this type for 1500 MW.

How much space would such a large number of turbines take? Here we have to account for the fact that a reasonable spacing is required. If wind turbines are too close, they will spoil each other's wind profile. This not only decreases the power of the wind turbines downstream, it also puts extra strain on the construction as a result of turbulence. It turns out that a spacing of 7 rotor diameters is a reasonable rule of thumb for wind farms. So the total area required is about 800 km². This is consistent with a rule of thumb saying that wind farms at sea generate, on average, between 1 and 2 MW per km², depending on type and location. This is, in first approximation, independent of the rotor diameter, since both turbine power and spacing scale with the square of the diameter. Large turbines obviously take advantage of the fact that the wind speed increases with altitude.

Given the size of the seas around Europe, 800 km² does not sound unreasonable. So we should opt for off-shore wind power?

Perhaps, but off-shore wind turbines have a drawback: building and maintaining them at sea is cumbersome, not to mention nasty corrosion by the sea water. This makes them roughly twice as expensive as turbines on land. Economically speaking, we would be better off with wind power on shore. Such turbines, if placed wisely, are almost comparable to traditional power plants. And their 'energy pay-back time' is less than a year. Sounds great, but it does not address our aesthetic objections.

One may wonder: How did our 17th century ancestors perceive the windmills that we find so charming in the Dutch landscape today? Interesting question. But the answer... is blowing in the wind.

30

Fog and raindrops



As every physicist knows, fog – or mist – is just a collection of tiny drops of water, at least if it is caused by nature. What distinguishes them from rain is, of course, their size. They are so small that their vertical speed is almost negligible. The dramatic effect of size on speed is obvious if we realize that for droplets smaller than, say, 0.1 mm, the flow profile around the droplet is purely laminar, so the friction F is determined by Stokes' law: $F = 6\pi\eta Rv$, with η the viscosity, R the radius and v the speed. And since the friction is balanced by weight, which is proportional to R^3 , we see that the speed is proportional to R^2 . This means that small droplets fall very slowly indeed. Take, for example,

water droplets of 2 μm diameter, much larger than the wavelength of light and therefore still visible. We find that they fall through air at a speed of about 0.1 mm per second. That's not particularly fast: even the slightest wind or air turbulence will offset such low speed.

But wait: do we really need turbulence to keep such tiny droplets airborne? Isn't thermal motion sufficient to keep them from falling? Don't they behave like ordinary molecules in the atmosphere, having a height distribution obeying Boltzmann's law? We can easily check if this is the case. We remember that Boltzmann's law implies a distribution over height h decaying as $\exp(-mgh/kT)$. In normal atmospheric conditions, the $1/e$ value is reached at a height of around 8000 m. Obviously, for particles much heavier than nitrogen or oxygen molecules we must settle for a distribution that stays closer to earth. Let us scale down the atmosphere for water droplet by a factor of one thousand, choosing a $1/e$ -value of 8 m. For this to be the case, the mass of a water droplet must be 1000 times that of a nitrogen or oxygen molecule, *i.e.*, it must consist of about 1500 water molecules. This is more like a large cluster than a droplet. Its diameter can be readily estimated by using the typical 'size' of 0.3 nm for small molecules or atoms in a liquid. In the case of water, we can even do a simple calculation if we consider a litre of water and use Avogadro's number. Sure enough, we find pretty exactly 0.3 nm for the distance between the centres of two neighbouring water molecules. From this it follows that the diameter of the cluster is only 5 nm. This is *really* small, much smaller than the wavelength of light. So we cannot see such clusters, but they surely make efficient light scatterers.

The conclusion? Mini-droplets smaller than about 5 nm would stay airborne forever, even in perfectly calm atmospheric conditions. They would form a perfect fog that never reaches the ground. If we were to walk or cycle through such a fog, it would be our *front* that got wet, not so much our head.

Alas, these mini-droplets do not survive very long. Inevitably, they collide and form larger drops. Slowly but surely they will start to fall. And by the time we can distinguish individual drops, we can be sure that we are walking in the rain.

31

Why planes fly



Ask any physicist how the wings of an aircraft work, and most probably he or she will come up with the popular explanation based on Bernoulli's law. The idea is that the cross section of a wing is curved along the upper side, and more or less flat at the bottom. Air hitting the front of the wing, the 'leading edge', is split in two, and the two air masses meet up again at the rear of the wing, the 'trailing edge'. Since the distance along the upper surface is longer, the air speed along the upper side must be greater. And according to Bernoulli's law, larger speeds imply lower pressures, and so there is a net upward force on the wing.

It sounds simple and logical. But it's wrong. We *know* it must be wrong. If this were the correct explanation, how on earth would planes be able to fly upside down?

So what *is* it that produces lift on a wing? It turns out that all we need is for the air flow to be deflected downward by the wing profile. As shown elegantly by Holger Babinsky from Cambridge back in 2003 (Physics Education **38**, p. 497-503), streamline curvature is the key. Think of a sailing boat, and forget the mast for a second. The sail can be seen as a vertical wing. It works beautifully propelling the boat, but its shape is nowhere near that of a traditional wing. There is no difference in path length along the two sides of the sail, so the Bernoulli explanation invoking different path lengths fails. Yet the sail is very efficient, simply because it creates curvature in the air flow. If we work it out, we find a simple relation between the curvature of the flow and the pressure gradient perpendicular to the streamlines: $dp/dn = \rho v^2/R$, with coordinate n normal to the streamlines, ρ the air density, v the speed and R the radius of curvature. The sign is such that the pressure decreases toward the center of curvature. This yields a pressure decrease at the convex side of the sail, and a pressure increase at the hollow side.

Indeed, thin curved wings like those of a sail are ideal for creating streamline curvature. Birds' wings tend to be like that. For aircraft, this is not an attractive option: thin curved wings would not meet structural demands and, in addition, would have no useful volume for storing fuel. Fortunately, *any* shape that introduces curvature into the flow profile can generate lift, even a symmetrical wing. All we have to do is to choose the 'angle of attack' appropriately: if the wing is slightly tilted upwards, its upper side will create streamline curvature as effectively as a thin curved wing would, thus giving by far the largest contribution to the lift. Below the wing there are regions of different senses of curvature, creating a net effect which is close to zero.

So, for a symmetrical wing the amount of lift – positive or negative – is purely a matter of adjusting the angle of attack, obviously within certain limits. And flying upside down is now a piece of cake. If you feel like it, of course.

32

Heating problems



It's winter; our house is warm and cozy, but we need fresh air. In an earlier daily-life column (*Fresh air*, page 25) we noticed that if we instantly refresh our room with cold outside air while keeping the heating off, the temperature will return almost to its original value once thermal equilibrium has been re-established. The reason is that the heat capacity of the air is small compared to that of all the solid-state stuff in our room, certainly if we include part of the walls. In turn, this is because the number density of atoms in solids is roughly 1000 times the density in air at ambient temperature and pressure, while the contribution to the heat capacity of every single

atom, whether in a solid, a liquid or a gas, is roughly the same: a few times the Boltzmann constant k .

Incidentally: just how instantaneous must the venting be in order for the argument to be valid? Obviously, the venting time has to be small compared to the thermal relaxation time of the room. But this is not simple to assess. For one thing: temperature equilibration of a room is not a single-relaxation-time process, given the wide variety in thermal relaxation times RC of all the objects in the room (R is the thermal resistance and C the heat capacity). The relaxation time of an empty wine glass, for example, may be just half a minute; the value for a full bottle of wine is about three hours, and that for the walls and other large objects may be even longer. In any case, the values are large enough so that ‘instantaneously’ refreshing the air is easily achieved. Subsequently, convection in combination with the small heat capacity of the air will rapidly raise the air temperature almost to its original level. Old-fashioned as this procedure may be, from an energy-saving perspective it has an advantage over having a continuous draught of cold air through our room, since this would make the temperature gradient near our skin steeper and make us feel cold.

The process of warming up the air in our room offers an interesting physics problem. If we compare the two situations: cold air and warm air in our room, in which of the two cases will the total kinetic energy of the air molecules in our room be largest, if we ignore convection currents? The answer seems obvious: since the mean kinetic energy of a gas molecule is directly proportional to temperature (*viz.*, $3/2 kT$), the total kinetic energy should go up.

But there is a catch. While the temperature goes up, some of the air will escape, since the atmospheric pressure will not change, being dictated by the outside pressure. And lo and behold, if we may consider air as an ideal gas (which is a very good approximation at ambient conditions) the density is inversely proportional to the temperature at constant pressure. And since the volume of our room remains constant, the number of molecules in it also decreases inversely proportional to the temperature.

So, the answer may be somewhat surprising: if we heat our room, the total kinetic energy of the air in it remains exactly constant.

33

Bubbles and balloons



When blowing soap bubbles as kids, we were probably much too fascinated by their beautiful colours to realise that there is some interesting physics going on. For one thing, the very existence of the bubbles demonstrates the concept of surface tension, since the slight overpressure inside the bubble has to be balanced by attractive forces in its 'skin'. And, in the process, it teaches us that, for a given volume, a sphere has the smallest surface area.

Blowing up a rubber balloon reveals some additional interesting aspects. Since the forces involved are much larger, some of these aspects are easily noticed. We have all experienced that the first stage

of blowing up a balloon is the hardest. Once the balloon has reached a certain volume, things get easier. The pressure needed decreases. This is funny, because everyone knows that, if you stretch a piece of rubber, the force required *increases* with length.

To fully understand the behaviour of the balloon, we need to know a bit more about the elasticity of rubber. This turns out to be significantly different to the normal behaviour of a common elastic material, for which Hooke's law holds: strain (relative change in length) is proportional to applied stress. For rubber, things are different. If we pull a piece of rubber band apart, we find that, after an initial rise in the stress similar to Hooke, there is a relatively flat plateau which ranges from strains of about 50% to 200%. Here the stress is reasonably constant. Only at about 400% – four times the initial length – does the stress increase steeply, since the macromolecules making up the rubber become fully stretched.

Now back to the balloon. Remembering the 'plateau' we assume for argument's sake that the 'surface tension' τ (force per unit length) is constant, just like in the case of soap bubbles. If we now consider a spherical balloon to consist of two imaginary halves and write down the force balance between the two halves ($\pi R^2 p = 2\pi R \tau$, with p the overpressure in the balloon), we find that the pressure needed to keep the balloon inflated is inversely proportional to the radius R . This qualitatively explains the fact that blowing-up the balloon gets easier once it has reached a certain size.

This observation calls for a spectacular experiment to amuse your audience. Take two balloons, inflate one of them to roughly one third of its maximum size and the other to two thirds. Attach both balloons to a piece of tubing while keeping the connection between the two closed with your finger. Ask the audience what will happen if you let go and connect both balloons through the tube. Sure enough, the audience expects the balloons to become equally big. After all, this is what happens if you take the two connected *uninflated* balloons and pull them apart: they will both be stretched to the same size.

But the audience is wrong. The big balloon gets bigger and the smaller one gets smaller. It illustrates the difference between force and pressure.

The funny properties of rubber are also at the heart of the remarkable behaviour which we see if we inflate a long, sausage-shaped balloon. We find that two ‘phases’ coexist at a single pressure. But here the physics is a bit more complicated. Not quite as simple as blowing bubbles.

34

Funny microwaves



The introduction of the microwave oven (or just ‘microwave’, the commonly used *pars pro toto*) has made our daily lives much easier. As physicists we may be a bit misled by the name and think of micrometer wavelength. But the standard oven operates at a frequency of 2.45 GHz, which corresponds to a wavelength of some 12 cm. That’s not precisely in the middle of the microwave region. This wavelength does explain though that, given the typical oven size, standing wave patterns can occur, causing large intensity differences over our food. We also realize that we need a convenient absorber in our food: water. But the absorption mechanism is not trivial. It is *not*

some intramolecular vibration or rotation mode that we are using. Typical rovibrational bands involve much higher energies, such that they are even responsible for the strong absorption in the red part of the visible spectrum in water. Instead, we use the large dipole moment of the water molecule to make it ‘wiggle’ amidst its neighbours. To be more precise: we absorb radiation by dielectric loss due to dipole relaxation. The microwave region is perfect for that. At much lower radiation frequencies the dipoles would follow the field changes and there would be no absorption. At very high frequency the dipoles have no time to change their orientation, and again nothing much happens. In between, where the dipoles lag behind the field, we expect a broad absorption curve. As already elucidated by Michael Vollmer in Physics Education back in 2004, the microwave frequency employed is not even near the maximum of the absorption curve. If that were the case, the absorption would be so large that only a thin layer of food would be heated. Instead, the frequency used is such that the penetration depth is in the order of a few cm, allowing our food to be heated more evenly.

An interesting consequence of the dipole relaxation mechanism is that ice has very little absorption. The molecules are simply too fixed in their lattice positions to follow the oscillating field. This reduces the absorption by three to four orders of magnitude.

So much for liquid or solid water in our oven: what about metals? Of course, reflection of the microwaves is almost perfect, due to the free electrons which essentially re-radiate the microwaves. Their penetration depth into the metal is in the order of 1 μm only. So our kitchen should be on perfectly safe ground as long as we keep the oven closed. And it should be no problem to leave a spoon in our cup of tea. A fork may be risky, though. Its sharp extremities will concentrate the electric field lines just like a lightning conductor does, and may lead to breakdown, with an interesting but possibly harmful light show as a result.

The most spectacular show may be caused by our precious decorated china, especially if the decoration is a thin gold layer. The reason is not trivial. We must remember the extremely small penetration depth of

the microwaves in metals. Inside that thin layer a lot of heat will be dissipated. For a solid metal piece like a spoon, this poses no problem. Its thermal conductivity and heat capacity are large, so it can easily absorb the heat and transfer it to the fluid in the cup. Alas, our china cups are poor thermal conductors, and the heat has nowhere to go but to the tiny thermal mass of the metal. So if we absent-mindedly put our beautifully decorated cup of tea in the microwave, we may have to kiss that cup goodbye...

35

Brave ducks



Remember how hard it was to first break the sound barrier? It took several fatal attempts by brave pilots before Charles ('Chuck') Yeager finally managed to fly faster than the speed of sound on the 14th of October, 1947. The problem was: by the time an aircraft approaches the speed of sound, the sound wave crests pile up in front of the plane. It then has to push through this barrier of compressed air in order to go faster than the waves. Once it is faster than the sound waves, an interesting situation occurs, quite similar to the case of a bullet moving at supersonic speed. The wave fronts produced have an enveloping circular cone, the 'Mach cone'. It is easy to see that the

half apex angle of the cone, θ , is related to the speed of sound c and the speed of the plane v by $\sin \theta = c/v$. Since there are no sound waves outside the Mach cone, the plane will pass us before we actually hear its sound.

Sound waves bear many analogies to water waves. Look at a duck, for example, speeding through a deep pond. See the V-shaped pattern of waves trailing the swimming duck? Doesn't it look like he is fighting the 'wave barrier' of water in front of him and producing a two-dimensional version of the Mach cone? Brave duck!

This certainly is an appealing thought. But it's wrong. What we may perceive as a 2-D version of a 'Mach cone' actually consists of two envelopes of a feathered pattern of dispersive waves.

Despite the analogies between water waves and sound waves, there are a few essential differences. Sound waves in air travel at a fixed speed without dispersion. The phase velocity c is equal for all wavelengths and equal to the group velocity. For supersonic flight this leads to the simple expression for the 'Mach angle' given above.

Water waves are much more complicated. They travel at the interface of two media, and are governed by gravity. Let us look at the deep-water limit, which is a good approximation for the duck as well as for ships on the ocean. Unlike sound waves in air, the phase velocity of the waves V depends on the wavelength, with long waves traveling faster than short waves. They follow the dispersion law $V = \sqrt{g/k}$ where g is the acceleration of gravity and k the wave number $2\pi/\lambda$. In other words, the speed of the waves is proportional to the square root of their wave length. At any speed of the duck or the ship, there will be waves running along with the same speed, whereas in the supersonic-flight case all waves are overtaken by the plane.

The complicated behaviour of the waves behind a duck or a ship in deep water was first worked out by Lord Kelvin (William Thomson), and is often referred to as 'Kelvin wake pattern' or 'Kelvin ship waves'. Kelvin was the first to find that, indeed, the wave pattern is bounded at either side by a straight line at an angle of 19.5 degrees with respect to the direction of the ship. This sounds like an awkward angle, and it results from a rather lengthy derivation. The angle may sound less

awkward if we write it down in its precise form, as $\arcsin(1/3)$. In turn, the $1/3$ results from the fact that the phase velocity given above is twice the group velocity. But the important thing is: this odd angle is fixed and characteristic for this type of wave. It has nothing to do with speed.

Too bad for the duck: In order to produce the V-shaped Kelvin wave pattern, he doesn't have to be brave and swim fast. Let alone faster than the 'speed of sound'.

36

Muddy cyclists



When watching the *Tour de France* or the *Giro d'Italia* on a not-so-sunny day, we are confronted with a simple physics problem. Why is it that cyclists on a wet road tend to get their back decorated with a vertical stripe of mud? Of course, it is due to the water from the road picked up by the tire. Centrifugal forces throw it off the tire somewhere in the upper part of its trajectory, and the forward speed launches it towards the poor cyclist's back. But why does the water leave the tire somewhere around the highest point? A superficial analysis of the wheel's motion may give us a clue. Any point along the rim traces out a cycloid, and its speed varies from zero to twice the speed of the bike. So isn't the answer

simply: it is because the speed of the tire rim is highest at its highest point, and so is the centrifugal force?

Reasonable as this may sound, it is entirely beside the point. Sure, it is the centrifugal force that counts. But that is the same everywhere along the wheel rim, given a certain speed. The fact that there is a linear motion superimposed on the wheel's rotation is irrelevant.

It is even quite the opposite, which we realize if we take gravity into account. Gravity tends to make the drops fall off much earlier, much closer to the road, whereas it tends to make the water stick to the tire near the top. We must conclude that the cyclist's back gets wet not *because of*, but *despite* the fact that the relevant tire part is near its highest position.

This raises the question: at exactly what speed does the cyclist get splattered with mud? We should realize that the drops leaving the tire precisely at its top position are rather innocent. They will leave horizontally, pass under the saddle and never make it to the cyclist's back. The real culprits are those drops that come off earlier, somewhere around 45 degrees before they reach the top, or even around 60 degrees before the top.

Now things get a bit complicated, since parameters like the exact position of the rider relative to the wheel come into play. Moreover, it is not sufficient to have centrifugal and gravitational forces balance. The water drops coming off the tire rim need some extra speed to be launched upward, in order to reach the cyclist's back.

A calculation for a standard cyclist and ignoring the drag on the droplets, done by Fokke Tuinstra from Delft University of Technology, shows that the drops which most likely make it to the rider's back will, indeed, leave the tire reasonably early, at around 60° before the top. They will hit the rider's back as soon as his speed exceeds some 12 km/h. If he rides a standard-size bicycle, that is. The reason is the crucial role of the wheel diameter. Given a certain speed v of the cyclist, the balance between centrifugal force and gravitation, $v^2/R = g$, shows that smaller wheels make things worse.

So, if you happen to be on your way to an important meeting wearing your business suit on a folding bike, you better make sure that the bike has an effective mud-guard over its back wheel.

37

Flying (s)low



When thinking about energy-efficient travel, why not use our imagination and try to construct a vehicle that has zero resistance? In fact, we do not have to invent it. It already exists. It's the airship, or zeppelin, named after its developer, the German count Ferdinand Graf von Zeppelin. It does not need high speeds to stay airborne, in contrast to a plane. Neither does it have the annoying rolling resistance of a car. So this looks like the ideal way of transport, viewed from a perspective of energy.

Or does it? All the above may be true if we just want the zeppelin to float at a fixed spot. But what happens to its efficiency once it starts moving?

We can easily make a back-of-an-envelope estimate. All we have to do is to work out the air resistance of the airship, keeping in mind that the resistance (in newtons) is equivalent to the energy dissipated per unit distance (in joules per metre, or kJ/km if you wish). To keep it simple, let us compare the airship with a car. This is a fair comparison: in contrast to a plane at high altitude, a car moves through air at ambient pressure, just like a zeppelin. After all, zeppelins are bound to fly low, since Archimedes' law would not allow them much lift in thin air.

And if we consider speeds of 100 km/h at the very least (just think of a zeppelin in a headwind!) the rolling resistance of the car can be ignored, since it makes only a minor contribution at such high speed.

So let us look at the air resistance, or drag. We may remember that it is given by $F = C_D A (\frac{1}{2} \rho v^2)$, where C_D is the drag coefficient, A the frontal surface area, ρ the air density and v the speed. For a fair comparison we should take the value of A *per passenger* in both cases. For a car, this is about 0.5 m². For a zeppelin we may take the dimensions of the Hindenburg, the airship that made history when it tried to land in New Jersey back in 1937. It had a diameter of 41 metres and carried about 100 passengers. This yields a frontal area of 13 m² per passenger. Obviously, there is no way that this can compete with a car. Even if we take into account that the value of C_D for the cigar-shaped zeppelin may be lower than the value for a car by a factor of three (0.1 vs. 0.3, say), the airship loses by an order of magnitude.

We can check our estimate using the Hindenburg's technical data. It had a top speed of 135 km/h and its engines had a power P of 3560 kW in total. If we work it out, realizing that $P = Fv$ we find that, indeed, a car beats the airship by a factor of 7 or 8.

If we remember that a full airplane is about half as fuel-efficient as a full car, we conclude that a plane is also superior to the zeppelin by a wide margin, even though its speed is much higher.

This may come as a surprise, but the reason is obvious. For one thing, the airship has this enormous volume, giving rise to large air resistance. Secondly, the density of the air through which it moves is larger by a factor of 4 compared with the air at cruising altitude of a plane.

The conclusion is inevitable. There is no bright future for the airship, even if the price of energy goes soaring. Unless we *really* take our time and go slow.

38

Funny ice



Water is a great substance, especially when it freezes. It becomes slippery and is fun to skate on. But why *is* it that ice is so slippery? It is not because it's flat. Glass, for example, is flat but not slippery. What we need is a layer of water to turn a flat surface into a slippery one. And sure enough, if we skate on ice we skate on a thin layer of water.

So where does the water come from? Many people think it's because of the pressure that the skater puts on the ice. After all, pressure lowers the melting point, and a skater's weight on a tiny skate makes quite a lot of pressure. But if you do the calculation

this turns out to change the melting temperature by a few tenths of a degree at most. So this explanation is wrong. Hardly surprising, when you consider that a hockey puck with its negligible weight slides so well across the ice.

In fact, we don't need pressure at all. There is always a thin layer of liquid water on the ice, up to some 70 nanometres thick if the temperature is just below freezing point. Basically this is because the molecules in the uppermost layer lack neighbours at one side, so they are not as tightly bound as the molecules in the bulk. Therefore ice is wet, and that allows us to glide so beautifully almost without any resistance.

So much for the skating fun. What about the freezing process, if we consider still water that is not flowing? Of course, we need sub-zero air temperatures to do the trick. And as long as the water temperature is above 4 °C, natural convection mixes the water, since warmer layers near the bottom are lighter and rise. But once the water is at 4 °C, it has reached its highest density, so the coldest water near the surface is lightest and remains on top. Convection stops, and the freezing process can begin. Since this situation is reached sooner in shallow than in deep water, this explains why shallow water freezes more easily.

If the air temperature rises again, is there something special about melting of the ice layer? For reasons of symmetry we may expect that the melting process is just as fast as the freezing, if the temperature differences are supposed to be equal and opposite.

Wrong. During freezing in still cold air, the air layer just above the ice is warmer than the rest of the air. Natural convection now helps to cool the ice. By contrast, if the ice is melting due to rising air temperature, the ice is relatively cold, so the cold air next to it will have no tendency to rise. Convection will *not* set in to increase heat transport. We conclude that melting is slower than freezing.

Skaters hate to see the ice disappear. Fortunately, as long as the air temperature remains sub-zero and if we can ignore radiation, the thickness of the ice layer should remain unchanged. Or does it? We realize that, even below 0 °C, there is a finite vapour pressure, so water molecules will go directly from the solid to the gaseous phase,

by sublimation. Many skaters will conclude that this is bad news, since it will decrease the thickness of the ice layer.

Wrong again. Sublimation cools the ice surface, the heat involved being the sum of melting and vaporization heat. This is almost an order of magnitude larger than just the melting heat, so the net effect is that this process makes the ice layer grow faster at the bottom than it disappears at the top.

So if you think that everything about water and its phase transitions is trivial, you're on thin ice...

39

Amazing candle flames



Granted: a candle flame is a lousy source of light. For an energy user of some 100 W, its light production doesn't even come close to any modern light source. But other than that, it represents an ingenious piece of technology.

Before going into details, we should realize that, when talking about flames, we are talking about chemical reactions in the *gas phase*. We can even illustrate this in a simple way by blowing out a candle, and re-lighting it by sticking a burning match into the stream of smoke above the hot wick. So, we cannot simply light a chunk of candle wax by a match, because its vapour pressure is far too low. Take

paraffin, by far the most common wax used in candle production. It consists of a mixture of hydrocarbons, for example C_nH_{2n+2} with n typically around 22 to 25. Such molecules have vapour pressures at ambient temperature far below 10^{-6} bar, much too low to be ignited and – fortunately – low enough for the candles to be stored almost indefinitely. So, for igniting paraffin we must get closer to its boiling point, which is somewhere in the range of 350 to 430 °C.

This is precisely what we achieve by having a wick with some paraffin absorbed. Its heat capacity is so small that its temperature can be raised by a burning match in just a second.

The wick is the heart of the candle. Its heat not only melts the wax just below it, it also acts as a fuel pump by drawing up liquid wax by capillary action, thereby regulating the flame. And if we look carefully after lighting a candle, we notice that the flame is large at first, then gets smaller by lack of fuel, and only burns in its full glory once it manages to melt a layer of wax. The wick is usually made of braided cotton threads, treated with some inorganic compound to prevent afterglow once the flame is extinguished. Its construction has a decisive impact on the performance of the candle, including the stance of the wick and its ability to self trim. It may contain a zinc or tin core to help it stay upright when the surrounding wax liquefies.

The operation of the candle teaches us, in passing, that the heat of combustion is much greater than the heat of melting and the heat of vaporization combined. It is one of the elementary physics lessons hidden in a candle.

Obviously, the operation of the candle depends on natural convection to remove the combustion products and supply fresh oxygen. Indeed, in microgravity a lit candle burns only for a short while before extinguishing – or *almost* extinguishing – by lack of oxygen. Here is another physics lesson: diffusion at ambient pressure is a very slow process.

The flame itself represents a series of steps: vaporization of the wax, pyrolysis into gaseous hydrocarbon fragments, hydrogen and solid carbon particles ('soot') and, finally, burning of the carbon particles in the luminous cone which is the whole purpose of the candle to begin with. In case of incomplete combustion of these C-particles

– for example, if there is lack of oxygen, or if a gust of wind decreases the flame's temperature to below $1000\text{ }^{\circ}\text{C}$ – the flame will emit soot and spoil the fun.

The temperature in the luminous cone is around $1200\text{ }^{\circ}\text{C}$. Now it becomes clear just why a candle is such an inefficient light source. Not only is more than 80% of the heat convected up and away from the flame. The remaining 20% does not provide very efficient lighting either. If we assume that the burning carbon particles at $1200\text{ }^{\circ}\text{C}$ behave like a Planck radiator, Wien's law tells us that its emission peak is at approximately $2\text{ }\mu\text{m}$ wavelength. Given the narrow eye sensitivity curve centered around $0.5\text{ }\mu\text{m}$, the conclusion is inevitable. Candles provide interesting science, but hardly any light.

40

Capricious sun-time



At what time of the day does the sun reach its highest point, or culmination point, when its position is exactly in the South? The answer to this question is not so trivial. For one thing, it depends on our location within our time zone. For Berlin, which is near the Eastern end of the Central European time zone, it may happen around noon, whereas in Paris it may be close to 1 p.m. (we ignore the daylight saving time which adds an extra hour in the summer).

But even for a fixed location, the time at which the sun reaches its culmination point varies throughout the year in a surprising way. In other words: a sundial, however accurately positioned, will show

capricious deviations through the seasons: the solar time on the sundial will almost always run slow or fast with respect to the ‘mean solar time’ on our watch. It’s all determined by the rotation of the earth around its axis, combined with its orbit around the sun.

The first thing we realise is that, from one day to the next, the earth needs to rotate a bit more than 360 degrees for us to see the sun in the South again. The reason is obvious. During a day, the earth moves a bit further in its orbit around the sun and thus needs to turn a little extra to bring the sun back to the same place (remember that the rotational direction of the earth around its axis *and* of its orbit around the sun are both counterclockwise). Now, if the earth were well-behaved, and would move in a circular orbit around the sun, with its rotational axis perpendicular to its orbital plane, this would be the end of the story.

But there are two complications, both of which cause deviations. The first one is the *elliptical* orbit of the earth. In fact, the earth is 3% closer to the sun at the beginning of January than at the beginning of July. So, the globe must rotate just a bit longer in January to have the sun back in the South than in July; just think of Kepler’s law. The result is that the solar time will gradually deviate from the time on our watch. We expect this ‘eccentricity effect’ to show a sine-like behaviour with a period of a year.

There is a second, even more important complication. It is due to the fact that the rotational axis of the earth is not perpendicular to the ecliptic, but is tilted by about 23.5 degrees. This is, after all, the cause of our seasons. To understand this ‘tilt effect’ we must realise that what matters for the deviation in time is the variation of the sun’s *horizontal* motion against the stellar background during the year. In mid-summer and mid-winter, when the sun reaches its highest and lowest point of the year, respectively, the solar motion is fully horizontal, so its effect on time is large. By contrast, in spring and autumn, the sun’s path also has a vertical component, which is irrelevant here. But it makes the horizontal component smaller in

these parts of the year, and so also its effect on time. This gives rise to a sine-like deviation having a period of *half* a year.

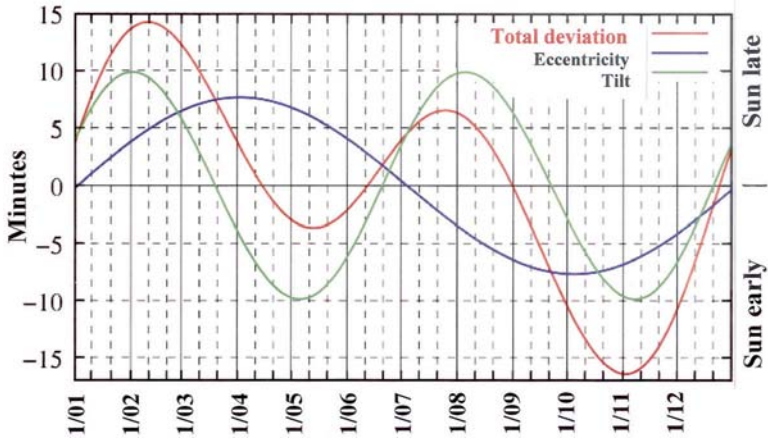


Figure 40.1 | Difference between solar time and ‘mean solar time’, and the separate contributions of the two underlying effects.

The two contributions are shown in the graph. Superposition of these ‘single and double frequency’ curves yields the total deviation of the ‘solar noon’ from the ‘mean solar noon’ on our watch. We see that around February 11 the sun is about 15 minutes later than average, and around November 3 about 15 minutes earlier.

So, a sundial in our front yard may be quite charming, but understanding its readings requires a scientist.

This page intentionally left blank